

MAXI FICHES

Électronique

Ludovic Barrandon

Maître de conférences à l'université de Picardie Jules Verne, INSSET
(Institut Supérieur des Sciences et Techniques)

Denis Réant

Professeur agrégé en électronique au lycée technique et professionnel
Condorcet, Saint-Quentin

Kambiz Arab Tehrani

Attaché d'enseignement et de recherche à l'université de Picardie Jules Verne,
INSSET

DUNOD



Le pictogramme qui figure ci-contre mérite une explication. Son objet est d'alerter le lecteur sur la menace que représente pour l'avenir de l'écrit, particulièrement dans le domaine de l'édition technique et universitaire, le développement massif du photocopillage.

Le Code de la propriété intellectuelle du 1^{er} juillet 1992 interdit en effet expressément la photocopie à usage collectif sans autorisation des ayants droit. Or, cette pratique s'est généralisée dans les établissements

d'enseignement supérieur, provoquant une baisse brutale des achats de livres et de revues, au point que la possibilité même pour

les auteurs de créer des œuvres nouvelles et de les faire éditer correctement est aujourd'hui menacée. Nous rappelons donc que toute reproduction, partielle ou totale, de la présente publication est interdite sans autorisation de l'auteur, de son éditeur ou du Centre français d'exploitation du droit de copie (CFC, 20, rue des Grands-Augustins, 75006 Paris).



© Dunod, Paris, 2010
ISBN 978-2-10-055476-8

Le Code de la propriété intellectuelle n'autorisant, aux termes de l'article L. 122-5, 2° et 3° a), d'une part, que les « copies ou reproductions strictement réservées à l'usage privé du copiste et non destinées à une utilisation collective » et, d'autre part, que les analyses et les courtes citations dans un but d'exemple et d'illustration, « toute représentation ou reproduction intégrale ou partielle faite sans le consentement de l'auteur ou de ses ayants droit ou ayants cause est illicite » (art. L. 122-4).

Cette représentation ou reproduction, par quelque procédé que ce soit, constituerait donc une contrefaçon sanctionnée par les articles L. 335-2 et suivants du Code de la propriété intellectuelle.

Table des matières

Table des matières	I
Avant-propos	1
LES BASES DE L'ÉLECTRONIQUE	
1 L'électron et l'électronique	2
2 Unités électriques et mesure	4
3 Électrostatique	8
4 Générateur, récepteur	10
5 Composants passifs	12
6 Valeurs des composants passifs	16
7 Notion d'impédance	18
8 Signaux périodiques	20
9 Utilisation des oscilloscopes	24
10 Lois de Kirchhoff	28
11 Loi d'Ohm	30
12 Énergie, puissance	32
OUTILS MATHÉMATIQUES ET PRINCIPES ÉLECTRONIQUES	
13 Notation et transformation complexes	34
14 Fonctions utiles en traitement du signal	36
15 Séries de Fourier	38
16 Transformée de Fourier	42
17 Transformée de Laplace	46

18	Le produit de convolution	48
19	Échelle logarithmique et décibels	50
20	Adaptation d'impédance	54
21	Quadripôles	56
22	Systèmes linéaires invariants dans le temps	60
MONTAGES À COMPOSANTS PASSIFS		
23	Montage parallèle, montage série	62
24	Théorèmes de Thévenin – Norton	66
25	Ponts de mesure	68
26	Filtres passifs du premier ordre	70
27	Filtres passifs du second ordre	74
SEMI-CONDUCTEURS		
28	Matériaux semi-conducteurs	78
29	Semi-conducteurs extrinsèques	80
30	Jonction PN et diodes	84
31	Différents types de diodes	88
TRANSISTORS BIPOLAIRES : PRINCIPES ET APPLICATIONS		
32	Transistor bipolaire	92
33	Réseau de caractéristiques	96
34	Polarisation d'un transistor	102
35	Modèle du transistor bipolaire	107
36	Amplificateur à transistor (principe)	112
37	Amplificateur émetteur commun	114
38	Montage collecteur commun	118

39	Amplificateur base commune	122
40	Amplificateur de classe A	124
41	Amplificateur de classe B	128
TRANSISTORS À EFFET DE CHAMP : PRINCIPES ET APPLICATIONS		
42	Transistors à effet de champ	132
43	Polarisation des TEC	136
44	Modèle du transistor à effet de champ	138
45	Montage à source commune	140
46	Montage à drain commun	144
AMPLIFICATEURS DIFFÉRENTIELS ET OPÉRATIONNELS		
47	Amplificateur différentiel	146
48	Amplificateur opérationnel	150
49	Montages à amplificateur opérationnel	152
50	Filtres actifs d'ordre 1 et 2	156
51	Synthèse de filtres actifs	158
TECHNOLOGIES NUMÉRIQUES		
52	Logique booléenne	162
53	Codes numériques	166
54	Portes logiques, logique combinatoire	170
55	Bascules, logique séquentielle	172
56	Circuits logiques programmables	174
57	FPGA	178
58	Le langage VHDL	182
59	Pratique du VHDL	186

60	Processeurs	191
INTERFACES ET COMMUNICATIONS		
61	Chaîne d'acquisition et échantillonnage	194
62	Convertisseurs Numérique-Analogique	198
63	Conversion Analogique-Numérique	200
64	Boucle à verrouillage de phase	204
65	Transmission de l'information	206
66	Modulations analogiques	208
67	Modulations numériques	212
68	Le bruit	216
ÉLECTRONIQUE DE PUISSANCE ET ALIMENTATIONS		
69	Introduction à l'électronique de puissance	218
70	Modèles simplifiés des semi-conducteurs de puissance	220
71	Composants semi-conducteurs de puissance	222
72	Convertisseurs statiques	226
73	Dimensionnement d'un dissipateur	232
ANNEXES		
74	Synthèse des trois montages à base de transistor bipolaire	238
75	Rappels mathématiques	240

Avant-propos

Cet ouvrage d'électronique, de la collection « maxi fiches », essentiellement destiné aux étudiants de niveau L1-L2 (SPI, EEA, Télécommunications...), BTS (Systèmes Électroniques, Contrôle Industriel et Régulation Automatique), IUT (GEII, Réseaux et Télécommunications, Mesures Physiques) saura également être un outil pratique pour les enseignants, les radioamateurs ou les techniciens et ingénieurs désireux de revenir sur certaines connaissances.

Son objectif est de faciliter l'acquisition et la révision des notions de base à travers 75 fiches synthétiques explorant les différentes parties du programme d'électronique des deux premières années du premier cycle universitaire et regroupées en 10 thèmes :

- Bases de l'électronique
- Outils mathématiques et principes électroniques
- Montages à composants passifs
- Semi-conducteurs
- Transistors bipolaires : principes et applications
- Transistors à effet de champ : principes et applications
- Amplificateurs différentiels et opérationnels
- Technologies numériques
- Interfaces et communications
- Électronique de puissance et alimentations

Les fiches sont généralement articulées en trois parties. La première, « En quelques mots » introduit et explique succinctement une notion, la seconde, « Ce qu'il faut retenir » en donne les définitions, les formules et les démonstrations essentielles et la troisième, « En pratique », permet au lecteur de voir concrètement les applications dans « le monde réel ». Les références croisées entre fiches s'apparentent à des liens hypertextes : elles facilitent un parcours non chronologique du livre et dénotent l'interdépendance entre les notions théoriques.

Nous espérons que ce recueil de fiches sera un allié précieux et pratique pour qui veut consolider efficacement ses connaissances théoriques en électronique. Nous souhaitons aux lecteurs autant de plaisir à utiliser cet ouvrage que nous en avons eu à le concevoir.

Ludovic Barrandon
Denis Réant
Kambiz Arab Tehrani

1 L'électron et l'électronique

Mots clés

Atome, orbite, charge électrique, bande de valence, bande de conduction, bandes d'énergie, semi-conducteurs.

1. EN QUELQUES MOTS

En toute généralité, les atomes sont constitués d'électrons tournant autour d'un noyau lui-même composé de neutrons et de protons. Un atome électriquement neutre comporte autant d'électrons que de protons. Les **électrons mis en mouvement coordonné** sont à la base de l'électricité, de l'électronique et, par extension, de l'électromagnétisme pour les radiocommunications.

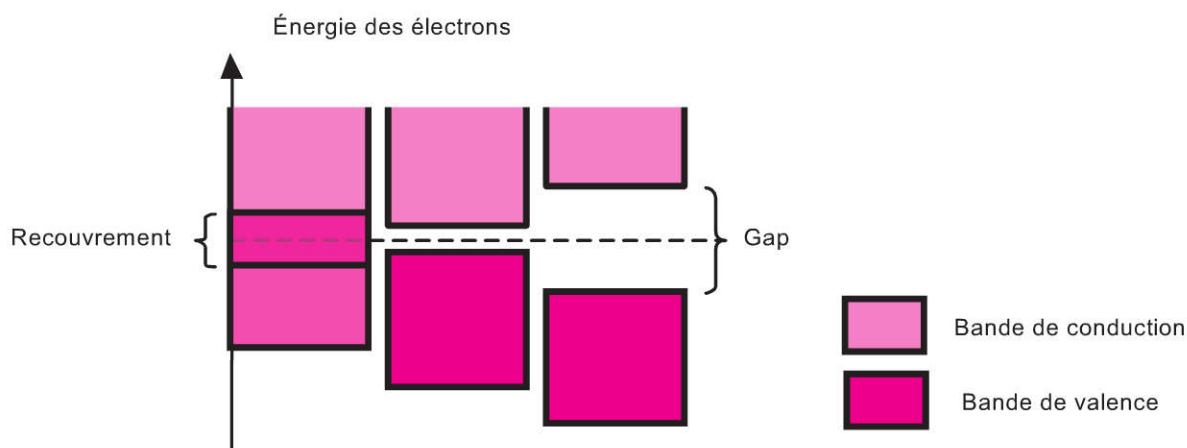
2. CE QU'IL FAUT RETENIR

En première approximation, la masse de l'électron, environ $1,67 \cdot 10^{-27}$ kg, est négligeable par rapport à la masse totale d'un atome. En revanche, sa **charge électrique, négative** par convention, est équivalente à celle d'un proton (fiche 2). Au sein d'un atome, les électrons se répartissent selon des orbitales atomiques.

Le remplissage des orbitales suit une logique bien déterminée. Il faut retenir qu'une orbitale correspond à un niveau d'énergie. Schématiquement, plus l'orbitale est éloignée du noyau et plus les électrons ont une énergie élevée, au même titre qu'une masse éloignée du sol présente de l'énergie potentielle. Les électrons les plus proches du noyau atomique sont liés à l'atome. Les électrons périphériques sont appelés électrons de valence (qui restent liés à l'atome) et **électrons de conduction** : ce sont ces derniers qui entrent en jeu dans le transport d'énergie électrique.

Dans un atome isolé, les orbitales correspondent à des niveaux d'énergie bien précis. Lorsqu'un grand nombre d'atomes sont réunis au sein d'un même solide, les orbitales deviennent très nombreuses et les niveaux d'énergie s'étalent pour former des **bandes d'énergie**.

Pour définir si un solide est un métal, un isolant ou un semi-conducteur, et donc s'il est susceptible ou non de conduire l'électricité, il faut s'intéresser à la différence d'énergie, appelée **gap**, qui existe entre la bande de valence et la bande de conduction. Si ces bandes se recouvrent, le solide est **conducteur**. Si elles sont très éloignées, le matériau est **isolant**. Le cas intermédiaire, celui des **semi-conducteurs**, très utilisés en électronique, présente des bandes de valence et de conduction proches mais distinctes (fiche 28).



Dans un métal, les électrons appartenant à la bande de conduction sont dits libres : ils peuvent se déplacer d'un atome à l'autre au sein du solide tout en transportant leur charge électrique. C'est le phénomène à l'origine de la conduction électrique.

3. EN PRATIQUE

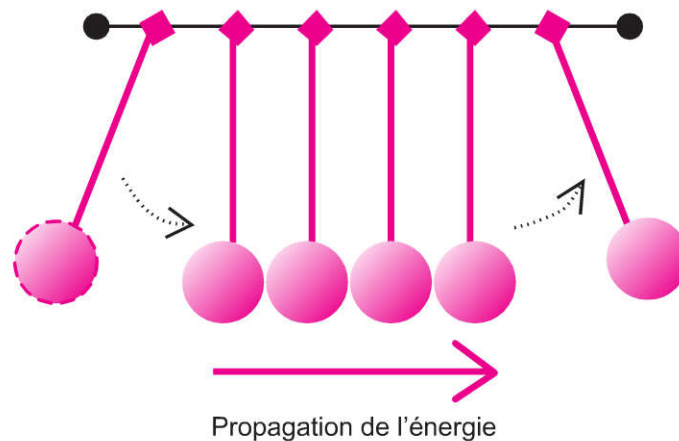
Les électrons se déplacent à une vitesse relativement lente : de l'ordre du centimètre par heure. En se référant à la [fiche 2](#), il est possible de mettre cette vitesse V en équation. Soit un conducteur de rayon r , parcouru par un courant continu I et dont la concentration en électrons libres est N .

$$V = \frac{I}{N \cdot q \cdot \pi \cdot r^2}$$

Le cuivre contient environ $8,5 \cdot 10^{28}$ électrons libres par mètre cube à température ambiante. Nous pouvons en déduire la vitesse des électrons dans un fil de cuivre de 2 mm de rayon parcouru par un courant de 1 ampère :

$$V = \frac{1}{8,5 \cdot 10^{28} \times 1,6 \cdot 10^{-19} \times \pi \times 4 \cdot 10^{-6}} \approx 0,00585 \cdot 10^{-3} \text{ m} \cdot \text{s}^{-1} \approx 2,1 \text{ cm} \cdot \text{h}^{-1}$$

Dans le cas d'un courant alternatif, les électrons se déplacent autour de leur position initiale. Comment se fait-il alors que la lumière du plafond s'allume instantanément lorsqu'on actionne l'interrupteur sur un mur situé à plusieurs mètres d'elle ? L'électricité se propage à une vitesse proche de celle de la lumière. Les choses se passent comme si les électrons se poussaient directement l'un l'autre dans le même sens, à l'image de billes métalliques suspendues à des fils : la bille de gauche en pointillés écartée de la position verticale redescend beaucoup moins vite vers sa position d'origine que l'énergie est transférée d'une bille à l'autre vers la dernière à droite.



En électricité, c'est au courant en tant que transport d'énergie que l'on s'intéresse. L'électronique quant à elle considère le courant électrique comme un **support d'information**.

2 Unités électriques et mesure

Mots clés

Unité, système international, ampère, volt, coulomb, watt, joule.

1. EN QUELQUES MOTS

Quel que soit le domaine physique considéré, une valeur numérique sans unité n'a pas de signification. Afin que tout le monde ait les mêmes repères, les unités ont été définies au niveau international. Le système international est également appelé MKSA puisqu'il est basé sur les unités fondamentales que sont le mètre, le kilogramme, la seconde et l'Ampère, dont découlent toutes les autres unités de la physique. Celles que nous rencontrerons principalement en électricité, et par extension en électronique, sont présentées ici.

2. CE QU'IL FAUT RETENIR

a) Unités

Les unités fondamentales (ou **unités de base**) sont définies à partir de phénomènes physiques. Ces phénomènes sont décrits de sorte qu'il n'y ait pas d'ambiguïté sur le résultat obtenu. Des unités fondamentales découlent toutes les autres unités : les **unités dérivées**.

Le premier exemple d'**unité fondamentale** du système international MKSA que nous allons définir est l'**Ampère**. C'est le Comité international des poids et mesures qui l'a défini en 1948.

Un Ampère est l'intensité d'un courant constant produisant une force égale à $2 \cdot 10^{-7}$ Newton par mètre linéaire, s'il est maintenu dans deux conducteurs linéaires et parallèles, de longueurs infinies, de sections négligeables, et distants d'un mètre dans le vide. Ceci est un exemple typique de **phénomène reproductible** utilisé comme base de définition d'une unité fondamentale.

Ce qu'il faut garder à l'esprit, c'est que l'intensité est un débit de charges. Si q est la charge inst-

stantanée qui circule entre deux points, l'intensité I peut se calculer comme suit : $I = \frac{dq}{dt}$

L'intensité est donc la dérivée en fonction du temps de la charge qui traverse une certaine section d'un matériau.

Le **Coulomb** est l'**unité d'une charge électrique** dans le système international. C'est une **unité dérivée** directement de l'Ampère car c'est la quantité d'électricité traversant une section d'un conducteur parcouru par un courant d'intensité d'un Ampère pendant une seconde ($1 \text{ C} = 1 \text{ s} \cdot \text{A}$).

La **charge électrique élémentaire** est, en physique, la charge d'un proton ou, de façon équivalente, l'opposé de la charge d'un électron. Elle est notée e . Un Coulomb est équivalent à $6,241\,509\,629\,152\,65 \cdot 10^{18}$ charges élémentaires. L'inverse de ce nombre nous donne la valeur en Coulomb d'une charge élémentaire e : elle vaut $1,602\,176\,487 \cdot 10^{-19} \text{ C}$.

Si on se réfère à la relation $I = dq/dt$, les Ampères sont donc équivalents à des Coulombs par seconde ($\text{C} \cdot \text{s}^{-1}$). Il suffit donc d'intégrer le courant par rapport au temps pour obtenir la charge correspondante :

$$dq = i(t) \cdot dt \Rightarrow q(t) = \int_0^t i(t) \cdot dt$$

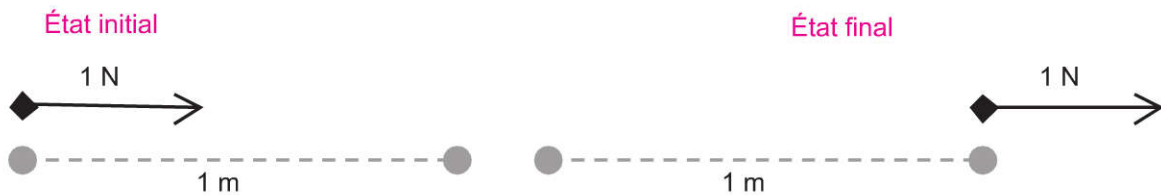
Si le courant est constant (courant continu), $Q = I \cdot \Delta t$ avec Δt l'intervalle de temps sur lequel on désire calculer la charge débitée.

Le **potentiel électrique** est souvent comparé à l'altitude d'un point. Le **Volt** est l'unité dérivée qui correspond à la **différence de potentiel électrique** qui existe entre deux points d'un circuit parcouru par un courant constant d'un Ampère lorsque la puissance dissipée entre ces deux points est égale à 1 Watt. Un Volt est donc équivalent à un Watt par Ampère ($\text{W} \cdot \text{A}^{-1}$). En utilisant les définitions du Coulomb et du Watt, on peut en déduire qu'un volt est également équivalent à un Joule par Coulomb car

$$V = \frac{W}{A} = \frac{J \cdot s^{-1}}{C \cdot s^{-1}} = \frac{J}{C}$$

La différence de potentiel électrique ou **tension électrique** correspond de manière imagée à une capacité à attirer les électrons.

L'**énergie** est une notion fondamentale de la physique. Ce concept est relativement difficile à appréhender. L'énergie se manifeste lorsqu'elle est dégagée ou reçue par un système : on peut observer des phénomènes tels que des mouvements, de la production de lumière, de chaleur... Quel que soit le phénomène physique considéré, l'énergie s'exprime en **Joules** dans le système international. L'énergie représente donc en quelque sorte la monnaie d'échange entre phénomènes physiques. Un Joule se définit comme le travail d'une force d'un Newton dont le point d'application se déplace d'un mètre dans la direction de la force, c'est donc une unité dérivée.



Entre l'état initial et l'état final, il aura donc fallu fournir une énergie d'un Joule. Dans le domaine de l'électricité, le Joule correspond à l'énergie fournie par le circuit électrique pour faire circuler un courant d'un Ampère pendant une seconde dans une **résistance** d'un Ohm. Ici, ce n'est pas tant le Joule que l'unité de résistance, l'**Ohm** (noté Ω), qui est définie. L'unité de résistance, l'Ohm, est donc dérivée du Joule, de l'Ampère et de la seconde. On peut déduire des définitions précédentes que

$$1 \text{ J} = 1 \text{ N} \cdot \text{m} = \text{A} \cdot \text{s} \cdot \Omega$$

La **puissance** est la capacité d'un système à dépenser ou à recevoir une certaine quantité d'énergie dans un intervalle de temps donné. Son unité, le **Watt** (noté **W**), correspond au transfert d'une énergie d'un Joule pendant une seconde :

$$W = J \cdot s^{-1} = V \cdot A^{-1}$$

b) Résumé

Grandeur physique	Unité		Signification	Équivalences
Tension	V	Volt	Différence de potentiel entre deux points	$V = \frac{W}{A} = \frac{J}{C}$
Courant	A	Ampère	Débit de charges	$A = C \cdot s^{-1}$
Charge	C	Coulomb	Propriété fondamentale de la matière	$C = A \cdot s$
Énergie	J	Joule	Force en action, capacité d'un système à produire un travail	$J = W \cdot s^{-1}$
Puissance	W	Watt	Dépense d'énergie	$W = J \cdot s$
Résistance	W	Ohm	Opposition au déplacement des électrons	$\Omega = V \cdot A^{-1}$

Afin de simplifier les écritures, toutes les unités peuvent être associées à un préfixe dénotant la puissance de 10 associée à la valeur indiquée : par exemple, 1 kV (un kiloVolt) représente 1 000 Volts au même titre qu'un kilogramme représente 1 000 grammes. Voici une liste non exhaustive des préfixes du système international :

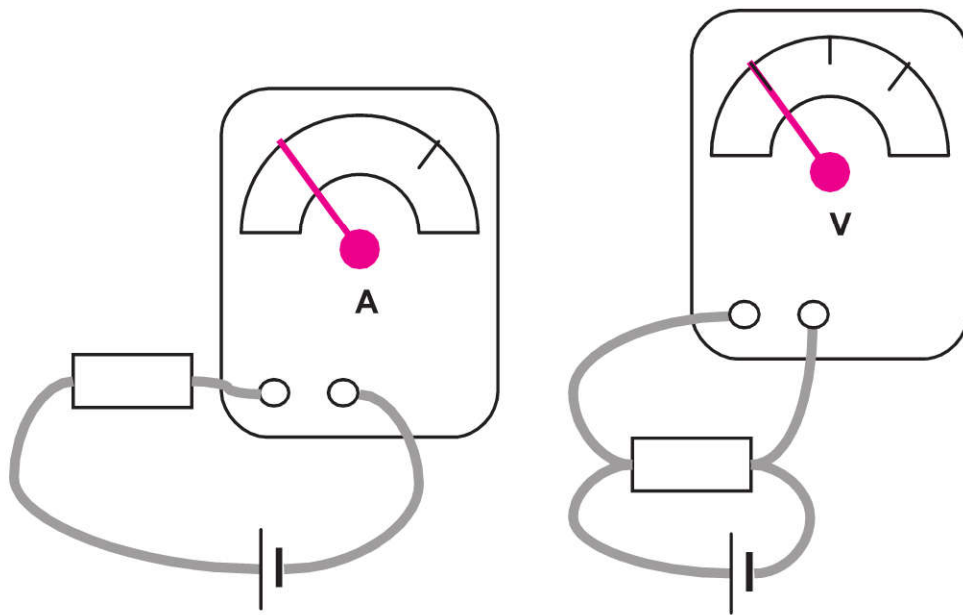
Femto	pico	nano	micro	milli		kilo	méga	giga	téra	péta
F	p	n	μ	m		k	M	G	T	P
10^{-15}	10^{-12}	10^{-9}	10^{-6}	10^{-3}	$10^0 = 1$	10^3	10^6	10^9	10^{12}	10^{15}

3. EN PRATIQUE

Différents appareils permettent de mesurer les grandeurs physiques dont il a été question dans le paragraphe précédent.

L'**ampèremètre**, comme son nom l'indique, affiche des valeurs en Ampère : il permet donc de mesurer des **courants**. Rappelons qu'un courant est un débit de charges. Afin de mesurer un débit, il convient d'insérer l'appareil en un point de la « canalisation » : un ampèremètre est inséré dans la maille du circuit dans laquelle on désire mesurer le courant. Pour ne pas perturber la mesure et le fonctionnement du circuit, l'appareil en lui-même ne doit pas s'opposer au passage du courant : sa résistance est négligeable. Le **galvanomètre** quant à lui est un type d'ampèremètre analogique très sensible.

Le **voltmètre**, comme son nom l'indique également, affiche des valeurs en Volt : cet appareil mesure donc des **tensions**. Une tension électrique étant une différence de potentiel, il faut brancher un voltmètre de part et d'autre de l'élément de circuit aux bornes duquel la tension est mesurée. Pour ne pas perturber la mesure et le fonctionnement du circuit, l'appareil en lui-même ne doit pas « détourner » les électrons de la portion de circuit étant l'objet de la mesure : idéalement sa résistance est infinie.



Les **multimètres** sont des appareils permettant entre autres la mesure de tensions et de courants : ils combinent donc les fonctions d'ampèremètre et de voltmètre. Ils offrent également souvent la possibilité de mesurer des résistances (ohmmètre) et parfois des capacités (capacimètre). Les modèles les plus répandus sont désormais à affichage numérique par opposition aux anciens modèles à aiguille.

L'**oscilloscope**, que l'on peut considérer comme un voltmètre un peu particulier, fait l'objet de la [fiche 9](#).

3 Électrostatique

Mots clés

Vecteurs, loi de Coulomb, force, charge électrique.

1. EN QUELQUES MOTS

L'électrostatique est l'étude des forces exercées par des charges électriques qui sont immobiles pour l'observateur.

2. CE QU'IL FAUT RETENIR

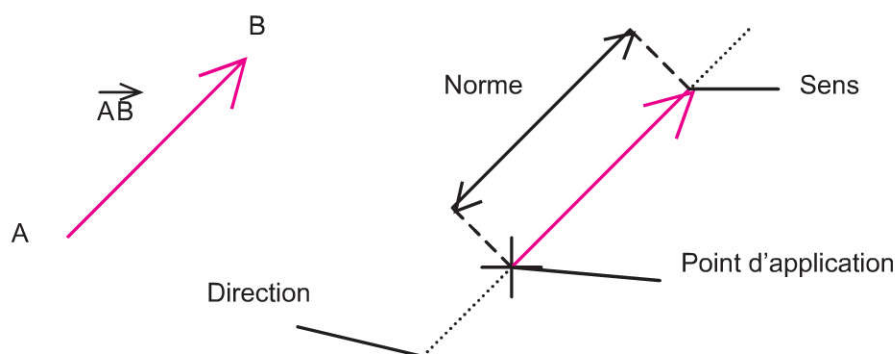
a) Rappels mathématiques

Une **quantité scalaire** est une quantité pouvant être décrite par un seul nombre et l'unité correspondante. Un phénomène scalaire est indépendant du sens d'un éventuel mouvement ou de sa direction. La masse, la température, la densité, les charges électriques, le potentiel électrique, la distance, le courant électrique, la tension, la puissance, sont des exemples de scalaires.

Les **vecteurs** sont extrêmement utiles pour décrire de manière précise le sens du mouvement et la direction de certains phénomènes. Un vecteur est défini par trois éléments :

- sa longueur également appelée norme ;
- sa direction ;
- son sens indiqué par une flèche.

Par exemple, pour illustrer une vitesse, une accélération, une force électrique, ou un champ électrique, etc. les indications de direction et de sens sont indispensables. Ces phénomènes sont vectoriels.



Propriétés :

- le produit de deux vecteurs est un scalaire ;
- le produit d'un vecteur et d'un scalaire est un vecteur.

b) Charge électrique

La charge électrique est une propriété fondamentale de la matière : c'est une notion abstraite qui permet d'expliquer certains phénomènes. Elle peut être positive ou négative. La force électrique s'exerce entre des particules électriquement chargées. À l'aide du concept de champ électrique, nous pouvons trouver la force électrique :

$$\vec{F}_e = |Q| \cdot \vec{E}_e$$

$\vec{F}_e$: Force électrique en Newtons appliquée sur la charge Q par le champ électrique.

Q : Charge électrique en Coulomb subissant la force électrique et située dans le champ.

$\vec{E}_e$: Champ électrique produit par les charges avoisinantes à la charge q ($\text{N} \cdot \text{C}^{-1}$)

Deux charges de signes opposés s'attirent et deux charges de même signe se repoussent, c'est ce que résume ce tableau :

Charge	+	-
+	Répulsion	Attraction
-	Attraction	Répulsion

c) Loi de Coulomb sous forme scalaire

La loi de Coulomb n'est valable que pour des charges immobiles ou très éloignées l'une de l'autre. Elle permet de calculer la force d'attraction ou de répulsion mutuelle créée par ces charges. Deux corps de charge q_1 et q_2 séparés par une distance d exercent l'un sur l'autre des forces soit attractives soit répulsives de même valeur F.

$$F = k_c \cdot \frac{q_1 \cdot q_2}{r^2}$$

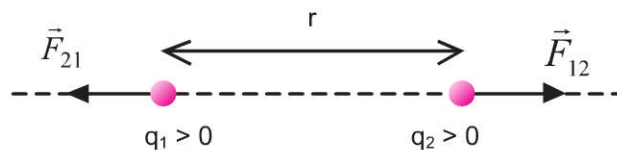
avec $k_c = 8,9875 \cdot 10^9 \cong 9 \cdot 10^9 \text{ N} \cdot \text{m}^2 \cdot \text{C}^{-2}$ dans le vide (ou dans l'air en première approximation). F s'exprime en N, q_1 et q_2 en C et d en m.

d) Loi de Coulomb sous forme vectorielle

Pour connaître la direction et le sens de la force provoquée par la présence de deux charges électriques, la forme vectorielle de la loi de Coulomb est nécessaire. Elle permet d'obtenir la grandeur et la direction de la force F_{12} exercée par une charge électrique q_1 placée au point de rayon vecteur r_1 sur une charge q_2 placée au point de rayon vecteur r_2 . Elle s'écrit :

$$|\vec{F}_{12}| = |\vec{F}_{21}| = k_c \cdot \frac{q_1 \cdot q_2 \cdot (r_1 - r_2)}{|r_1 - r_2|^3} = k_c \cdot \frac{q_1 \cdot q_2}{r^2} \cdot \hat{r}_{21}$$

Les forces $\vec{F}_{21}$ et $\vec{F}_{12}$ sont de même norme, c'est-à-dire que l'action des deux charges est réciproque.



La force totale sur une charge q (immobile) est la somme vectorielle des forces exercées par chacune des autres charges qui l'entourent.

$$\vec{F}_T = \vec{F}_1 + \vec{F}_2 + \vec{F}_3 + \dots + \vec{F}_{n-1} + \vec{F}_n$$

3. EN PRATIQUE

Les forces électriques sont à l'origine du mouvement des électrons dans les circuits : elles permettent de représenter sous forme mathématique l'attraction des électrons de la borne + vers la borne - d'un générateur. Un champ électrique apparaît lorsque deux charges électriques sont éloignées l'une de l'autre : par exemple, ce phénomène apparaît dans les condensateurs plans (fiche 5) et dans les diodes à jonction (fiche 30). Lorsque les charges sont en mouvement, la mise en équation est plus complexe et sort du cadre de cet ouvrage. Nous pouvons signaler que l'électrodynamique classique décrit le mouvement des particules chargées : nous rentrons alors dans le domaine de l'électromagnétisme.

4 Générateur, récepteur

Mots clés

Force électromotrice, énergie, loi d'Ohm, résistance interne, générateur réel, générateur idéal.

1. EN QUELQUES MOTS

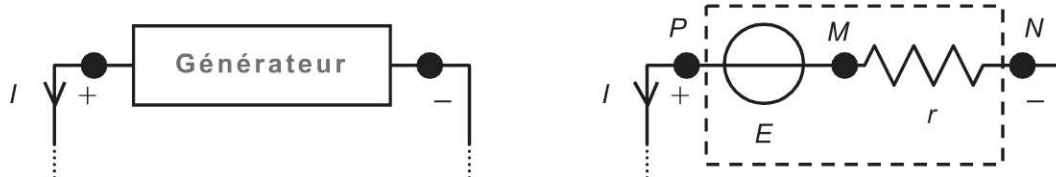
Un générateur électrique produit de l'énergie **à partir d'une autre forme d'énergie**, qu'elle soit mécanique (dynamo, éolienne), chimique (pile, batterie), photovoltaïque (panneaux solaires) ou thermique (centrale nucléaire). Le générateur fournit de l'électricité à un ou plusieurs **récepteurs** dans lesquels elle est utilisée.

2. CE QU'IL FAUT RETENIR

a) Générateurs

Un **générateur de tension** réel peut être représenté par un générateur de tension idéal en série avec une **résistance interne**. Un générateur de tension idéal est un dipôle pouvant fournir une tension constante quelle que soit la charge branchée à ses bornes. La **force électromotrice** d'un générateur est la tension que celui-ci peut délivrer au réseau.

La figure ci-après distingue un générateur de Thévenin d'un générateur réel de tension.



L'énergie W_T dépensée par le générateur est la somme de l'énergie W_I fournie au réseau (courant I) et de l'énergie W_2 dissipée par la résistance interne.

$$W_1 = Q \cdot (V_P - V_N) = I \cdot t \cdot (V_P - V_N) \quad W_2 = r \cdot I^2 \cdot t$$

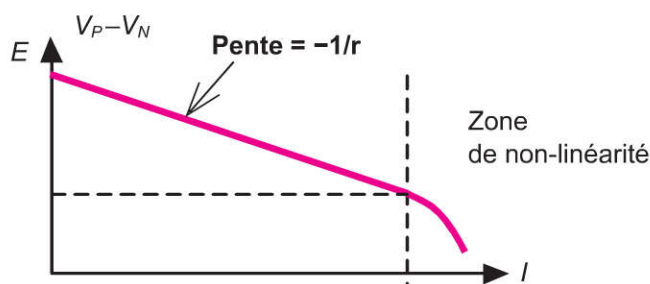
$$W_T = W_1 + W_2 = I \cdot t \cdot (V_P - V_N) + r \cdot I^2 \cdot t$$

La puissance totale P_T délivrée par le générateur peut se calculer en intégrant l'énergie W_I par rapport au temps.

$$\left. \begin{array}{l} P_T = (V_P - V_N) \cdot I + r \cdot I^2 \\ P_T = E \cdot I \end{array} \right\} \Rightarrow E = (V_P - V_N) + r \cdot I$$

On en déduit la différence de potentiel $V_P - V_N$ aux bornes du générateur de tension réel

$$V_P - V_N = E - r \cdot I$$

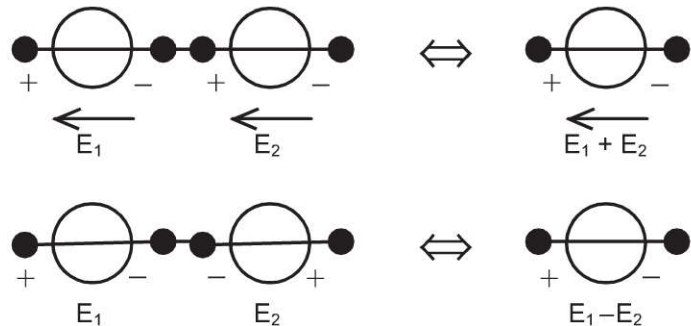


Le rendement en puissance ρ du générateur est le rapport entre la puissance délivrée par le générateur idéal de tension et la puissance délivrée au réseau seul : la puissance dissipée par la résistance interne est donc considérée comme perdue (fiche 12).

$$\rho = \frac{(V_P - V_N) \cdot I}{E \cdot I} = 1 - \frac{r \cdot I}{E}$$

Remarques

- Il ne faut pas placer deux générateurs de tension en parallèle : le générateur délivrant la tension la plus forte débitera du courant dans le générateur de tension le plus faible.
- Les tensions de deux générateurs de tension placés en série s'additionnent (loi des mailles).

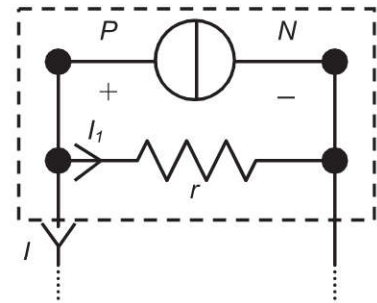


Un **générateur de courant réel** peut

être représenté par un générateur de courant idéal en parallèle avec sa résistance interne. Un générateur de courant idéal est un dipôle pouvant fournir un courant constant quelle que soit la charge branchée à ses bornes. La figure ci-après représente le générateur de Norton d'un générateur réel de courant.

De la même façon que pour le générateur de tension réel,

$$I = \frac{E - (V_P - V_N)}{r}$$



Remarque

Les courants de deux générateurs de courant placés en parallèle s'additionnent (loi des nœuds).

b) Application de la loi d'Ohm (fiche 11)

Pour calculer la différence de potentiel entre deux points d'un circuit, il suffit de se déplacer entre ces deux points.

- Si on rencontre une résistance R , on compte une différence de potentiel de :
 - $+R \cdot I$ si le déplacement se fait dans le sens du courant ;
 - $-R \cdot I$ si le déplacement se fait dans le sens inverse du courant.
- Si on rencontre un générateur de force électromotrice E ou un récepteur de force contre-électromotrice E , on compte :
 - $+E$ si on rentre par la borne positive ;
 - $-E$ si on rentre par la borne négative.

c) Source idéale indépendante

Une source de courant peut être considérée comme un circuit ouvert lorsqu'elle est éteinte ou qu'elle fonctionne à une fréquence ω différente de la fréquence ω_0 d'analyse du système.

Une source de tension peut être considérée comme un court-circuit lorsqu'elle est éteinte ou qu'elle fonctionne à une fréquence ω différente de la fréquence ω_0 d'analyse du système.

5 Composants passifs

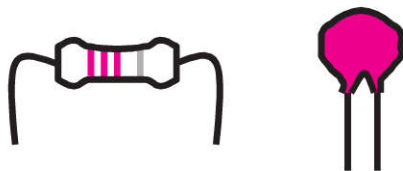
Mots clés

Résistance, ohm, condensateur, farad, inductance, henry.

1. EN QUELQUES MOTS

Les composants passifs ne nécessitent pas d'alimentation spécifique pour fonctionner. Les résistances, condensateurs, inductances, transformateurs et diodes (fiches 30 et 31) et tout assemblage de ces composants (filtres passifs (fiches 26 et 27), ponts de diodes...) rentrent dans cette catégorie.

Dans le domaine des faibles puissances propres à l'électronique, nous trouvons deux grandes techniques de fabrication : les composants à broches ou « traversants » et les composants montés en surface (CMS).



Les CMS sont brasés à la surface du circuit imprimé : ils ne nécessitent donc pas de perçage et sont globalement plus compacts.

473

2. CE QU'IL FAUT RETENIR

a) Résistances

Les résistances s'opposent au passage du courant : elles répondent à la forme élémentaire de la loi d'Ohm (fiche 11). Voici leur symbole européen (à gauche) et américain (à droite).



Les résistances sont couramment fabriquées à partir d'un dépôt métallique sur support isolant. La valeur d'une résistance dépend de la résistivité du matériau et des dimensions géométriques du composant.

$$R = \frac{\rho \cdot l}{S} = \frac{1}{\gamma \cdot S}$$

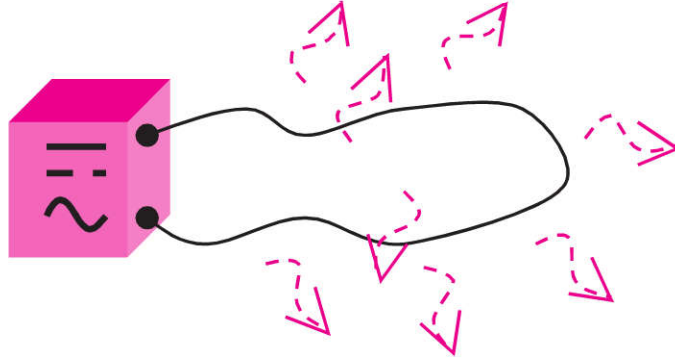
R est la résistance en ohms, ρ la résistivité en $\Omega \cdot m$, γ la conductivité en $\Omega^{-1} \cdot m^{-1}$, l la longueur du composant en m et S la surface de la section en m.

Certaines résistances sont variables : l'une de ses deux bornes peut se déplacer le long de la partie résistive, ce sont des potentiomètres.

Si l'on place un conducteur aux bornes d'un générateur, l'énergie sera principalement dissipée sous forme de chaleur, c'est ce que l'on appelle l'**effet Joule**. Comme nous l'avons vu, tout conducteur a

une résistance qui dépend du matériau dont il est constitué et de ses dimensions géométriques. La puissance dissipée P (en Watt) par effet Joule est $P = R \cdot I^2$.

I étant l'intensité du courant, en ampères, traversant la résistance et R la valeur de la Résistance, en Ohm. Pour les courants alternatifs, il convient de remplacer I par la valeur efficace I_{eff} du courant (fiche 8).

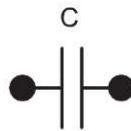


En-deçà d'une certaine température dépendante de chaque matériau, la résistivité du matériau devient nulle (il n'y a donc pas de dissipation d'énergie par effet Joule) : c'est la **supraconductivité**. Voici quelques exemples de températures, en Kelvin, pour lesquelles certains métaux n'opposent plus de résistance au passage du courant :

- plomb (Pb) : 7,2 K
- étain (Sn) : 3,8 K
- aluminium (Al) : 1,15 K

b) Condensateurs

Les condensateurs sont des composants de stockage de charges électriques.



Ils sont fabriqués à partir de deux plaques de conducteurs ou armatures schématisées sur le symbole normalisé, de surface S , placées face-à-face et séparées d'une distance d . Le volume situé entre les deux plaques est rempli d'un diélectrique de permittivité ϵ . La capacité du con-

densateur répond à la relation suivante : $C = \frac{\epsilon \cdot S}{d}$

La capacité s'exprime en Farad si S est en m^2 , d en m et ϵ en $\text{F} \cdot \text{m}^{-1}$.

Certains condensateurs sont munis d'armatures mobiles qui permettent de modifier la capacité du condensateur : ce sont les varicaps (fiche 31).

Lorsqu'une différence de potentiel V est appliquée aux bornes d'un condensateur de capacité C , ce dernier contient une charge q telle que $q = C \cdot V$

Sachant que $i(t) = \frac{dq}{dt}$ (fiche 2), il en découle que $i(t) = C \cdot \frac{dv(t)}{dt}$

L'impédance complexe d'une capacité est purement imaginaire :

$$\underline{Z}_C = \frac{1}{j \cdot C \cdot \omega} = -\frac{j}{C \cdot \omega} \text{ avec } \underline{U} = \underline{Z}_C \cdot \underline{I}$$

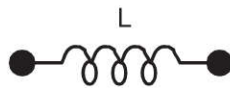
Les condensateurs peuvent éventuellement être polarisés, c'est-à-dire qu'ils présentent une borne + qui doit être reliée au potentiel électrique le plus fort et une borne – reliée au potentiel électrique le plus faible. Une erreur d'inversion peut entraîner une destruction, éventuellement explosive, du composant. Les condensateurs sont soit de type électrolytique (ou chimique dans le langage courant), soit de type tantale.

Les condensateurs non polarisés, de type céramique ou mylar, n'ont pas de sens de branchement mais ils ne permettent pas d'atteindre des valeurs aussi élevées que les condensateurs polarisés.

c) Inductances

Les déplacements des électrons engendrent un champ électromagnétique qui a tendance à s'opposer au déplacement relatif des électrons dans les fils voisins. De proche en proche, le courant électrique est d'autant plus affaibli qu'il y a de boucles (ou spires), que ces boucles sont resserrées ou que le courant varie rapidement.

Le symbole normalisé représente cette constitution en spires.



Les inductances sont des composants constitués d'un fil conducteur spiralé sur une ou plusieurs couches. Les fortes valeurs d'inductance sont atteintes en ajoutant un noyau magnétique à l'intérieur du cylindre ainsi formé.

Le Henry est l'unité d'inductance : $H = V \cdot A^{-1} \cdot s$

La valeur L d'une inductance à une seule couche de spires s'exprime de la manière suivante :

$$L = \frac{\mu_0 \cdot \mu_r \cdot N^2 \cdot S}{l}$$

L = inductance en henry (H)

μ_0 = constante magnétique = $4\pi \cdot 10^{-7} H \cdot m^{-1}$

μ_r = perméabilité relative effective du matériau magnétique. Elle est de 1 pour les bobines à air.

N = nombre de spires

S = section de la bobine en mètres carrés (m^2)

l = longueur de la bobine en mètres (m)

Le condensateur correspond à lier temporellement le courant à une dérivée de la tension. Le comportement électrique de l'inductance est réciproque puisqu'il relie la tension à une dérivée

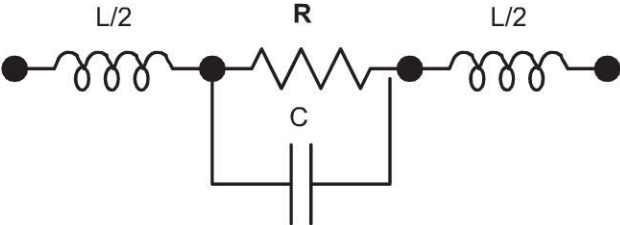
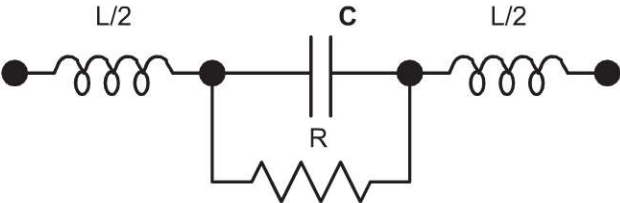
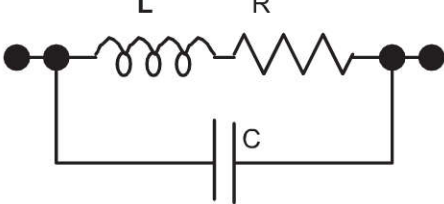
du courant : $u(t) = L \cdot \frac{di(t)}{dt}$

Traduite dans le domaine complexe, l'impédance s'écrit alors $\underline{Z}_L = j \cdot L \cdot \omega$ avec $\underline{U} = \underline{Z}_L \cdot \underline{I}$

3. EN PRATIQUE

Utilisés seuls, les composants passifs sont employés pour des fonctions d'atténuation du signal ou des alimentations électriques des composants (polarisation, [fiches 34 et 43](#)). Cette atténuation est soit indépendante de la fréquence du signal d'entrée, ce qui est le cas pour les montages diviseur de tension et diviseur de courant ([fiche 23](#)), soit variable en fonction de la fréquence d'entrée dans le cas des filtres ([fiches 26 et 27](#)).

Le comportement des composants étudiés ici n'est pas parfait. Lorsque la fréquence d'utilisation augmente, des comportements parasites apparaissent. Voici les schémas équivalents des trois composants de cette fiche pour les hautes fréquences.

Composant	Schéma équivalent hautes fréquences	Impédance équivalente	
Résistance		$f \ll f_R$	R
		$f \approx f_R$	$L \cdot \sqrt{\frac{1}{R^2 \cdot C^2 + L \cdot C}}$
		$f \gg f_R$	$j \cdot L \cdot \omega$
Condensateur		$f \ll f_R$	$-j \cdot \frac{1}{C} \cdot \omega$
		$f \approx f_R$	$L \cdot \sqrt{\frac{1}{R^2 \cdot C^2 + L \cdot C}}$
		$f \gg f_R$	$j \cdot L \cdot \omega$
Inductance		$f \ll f_R$	$j \cdot L \cdot \omega$
		$f \approx f_R$	$\frac{1}{R \cdot C \cdot \omega} \cdot \sqrt{R^2 + L^2 \cdot \omega^2}$
		$f \gg f_R$	$\frac{1}{C \cdot \omega} \cdot \sqrt{\frac{L^2}{R^2 + L^2}}$

6 Valeurs des composants passifs

Mots clés

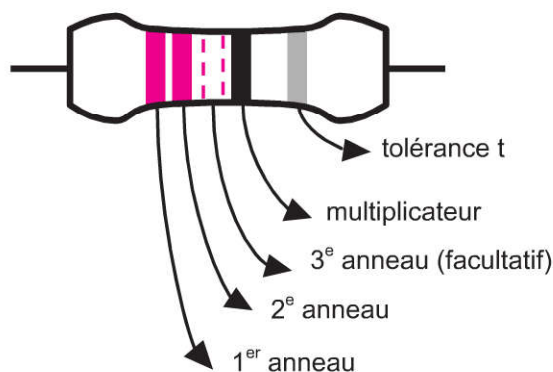
Résistance, capacité, inductance, valeurs normalisées, tolérance.

1. EN QUELQUES MOTS

Les composants électroniques sont généralement de petites dimensions. Les composants passifs (*fiche 5*) tels que les résistances, condensateurs et inductances ont des dimensions de l'ordre de quelques millimètres (composants traversants) voire moins pour les composants montés en surface (CMS). Il est donc difficile d'indiquer la valeur exacte d'un composant sur son propre boîtier en incluant l'unité et la tolérance associée, c'est pourquoi on a recours à des techniques codant les valeurs et les rendant plus lisibles.

2. CE QU'IL FAUT RETENIR

Classiquement, les **valeurs de résistances** sont indiquées directement sur le composant. Un des moyens utilisés pour les résistances est le code de couleurs standard issu de la norme internationale CEI 60 757. Si la résistance comporte quatre anneaux, les deux premiers correspondent chacun à une valeur numérique comprise entre 0 et 9, le troisième à un coefficient multiplicateur (puissance de 10) et le quatrième à la tolérance, c'est-à-dire à l'erreur relative permise sur la valeur nominale. Si la résistance comporte cinq anneaux, les trois premiers correspondent à une valeur numérique, les deux derniers ont la même fonction qu'avec quatre anneaux.



Dans le cas où la résistance comporte quatre anneaux, la valeur sera déduite avec le calcul suivant :

$$R = (10 \times N_{\text{anneau 1}} + N_{\text{anneau 2}}) \cdot 10^{N_{\text{anneau 3}}} \pm R \times t$$

$$\text{Avec cinq anneaux, } R = (100 \times N_{\text{anneau 1}} + 10 \times N_{\text{anneau 2}} + N_{\text{anneau 3}}) \cdot 10^{N_{\text{anneau 4}}} \pm R \times t$$

Il existe différentes phrases qui sont autant de moyens mnémotechniques pour retenir l'ordre des couleurs du tableau ci-dessous :

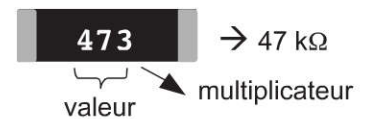
Noir	Marron	Rouge	Orange	Jaune	Vert	Bleu	Violet	Gris	Blanc
Ne	Manger	Rien	Ou	Jeûner	Voilà	Bien	Votre	Grande	Bêtise

La limitation de ce type de phrases est la répétition de deux initiales : le v et le b. Pour lever l'ambiguïté, on peut se rappeler intuitivement que le blanc (9) est à l'opposé du noir (0) et que le vert que l'on peut obtenir par mélange du bleu et du jaune se situe justement entre ces deux couleurs.

Couleur	1 ^{er} anneau	2 ^e anneau	3 ^e anneau	Multiplicateur	Tolérance	Série
Noir	0	0	0	$10^0=1\ \Omega$		
Marron	1	1	1	$10^1=10\ \Omega$	+/- 1 %	E96
Rouge	2	2	2	$10^2=100\ \Omega$	+/- 2 %	E48
Orange	3	3	3	$10^3=1\ \text{k}\Omega$		
Jaune	4	4	4	$10^4=10\ \text{k}\Omega$		
Vert	5	5	5	$10^5=100\ \text{k}\Omega$	+/- 0,5 %	E192
Bleu	6	6	6	$10^6=1\ \text{M}\Omega$	+/- 0,25 %	
Violet	7	7	7	$10^7=10\ \text{M}\Omega$	+/- 0,10 %	
Gris	8	8	8	$10^8=10\ \text{M}\Omega$	+/- 0,05 %	
Blanc	9	9	9	$10^9=1\ \text{G}\Omega$		
Or				$10^{-1}=0,1\ \Omega$	+/- 5 %	E24
Argent				$10^{-2}=0,01\ \Omega$	+/- 10 %	E12

La valeur de résistance est également souvent écrite sous forme chiffrée sur le composant : c'est le cas pour les **CMS** ou **composants montés en surface**. Le principe est similaire à celui du code couleur. Si le code comprend quatre chiffres, les trois premiers représentent la valeur et le quatrième le multiplicateur. Si la valeur est inférieure à 100, la lettre R indique la position du séparateur décimal.

Le principe est le même avec le marquage à trois chiffres, seuls les deux premiers représentent la valeur et le troisième est le multiplicateur. La lettre est utilisée comme séparateur pour les valeurs inférieures à 10.



La valeur des condensateurs est parfois écrite sur le boîtier en utilisant le code numérique décrit pour les résistances mais en se référant non pas à des Farad mais à des pF ($10^{-12}\ \text{F}$). De même, la valeur des inductances est parfois écrite sur le boîtier en utilisant le même code numérique en se référant à des μH ($10^{-6}\ \text{H}$).

3. EN PRATIQUE

Les valeurs normalisées de résistances, et dans le cas général des composants passifs, sont réparties de telle sorte que deux résistances de valeurs nominales consécutives ne puissent se recouvrir en tenant compte de la tolérance. Soient R_1 une valeur de résistance normalisée, R_2 la résistance de valeur nominale suivante et t la tolérance de la série E_p . La valeur p désigne le nombre de valeurs normalisées par décade. Soit α le coefficient multiplicateur entre la valeur de résistance R_1 et la valeur R_2 suivante, on a : $R_2 = \alpha \cdot R_1$. Les composants étant normalisés par décade, cela implique que

$$\left. \begin{array}{l} R_{n+p} = 10 \times R_n \\ R_{n+p} = \alpha^n \times R_n \end{array} \right\} \Rightarrow \alpha = \sqrt[n]{10}$$

Voici les différentes valeurs normalisées pour une décade dans la série E12.

n	0	1	2	3	4	5	6	7	8	9	10	11	12
$(\sqrt[12]{10})^n \approx$	1,00	1,21	1,47	1,78	2,15	2,61	3,16	3,83	4,64	5,62	6,82	8,25	10,0
Arrondi à	1,0	1,2	1,5	1,8	2,2	2,7	3,3	3,9	4,7	5,6	6,8	8,2	10

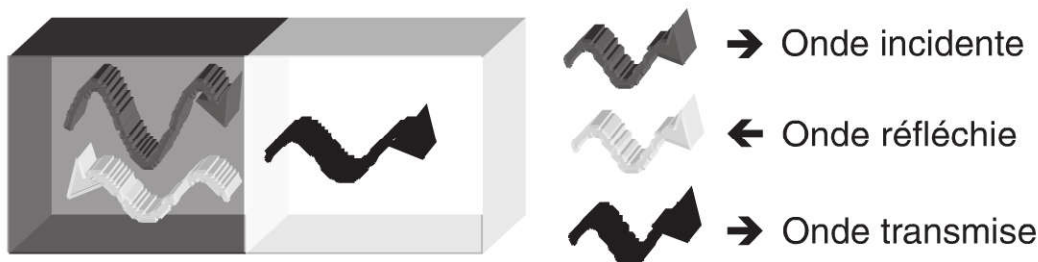
7 Notion d'impédance

Mots clés

Ondes incidentes, réfléchies, transmises, réactance, résistance, conductance, susceptance, admittance.

1. EN QUELQUES MOTS

La notion d'impédance n'est pas propre à l'électricité et à l'électronique. Elle se retrouve également dans différents domaines ayant trait aux forces mécaniques : acoustique, hydraulique... Lorsqu'une **onde se propage** d'un milieu vers un autre dont les caractéristiques sont différentes, une partie de l'**énergie est réfléchi**e avec une **certaine atténuation et un déphasage en fonction de la fréquence** et l'autre est transmise dans le milieu également avec un déphasage et de l'atténuation. C'est ce qui se passe lorsqu'on parle dans une pièce : une partie de la voix est entendue dans la pièce voisine, plus faiblement, et une autre est réverbérée dans la pièce elle-même.



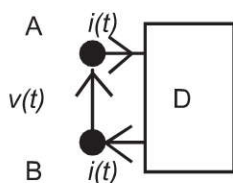
Dans le domaine de l'électronique, cette notion est une **généralisation de la loi d'Ohm pour les courants alternatifs**.

2. CE QU'IL FAUT RETENIR

Par définition, l'impédance rend compte de l'opposition d'un circuit au passage d'une onde électrique sinusoïdale.

Le formalisme utilisant des impédances complexes (fiche 13) et qui est une généralisation de la loi d'Ohm écrite en premier lieu pour les courants continus (fiche 11) n'est valable que pour des composants linéaires, c'est-à-dire modélisables sous forme d'équations linéaires. Pour les self-inductances (L) et les capacités (C), l'impédance et son inverse, l'**admittance**, sont fonctions de la fréquence. Dans tous les cas, un signal périodique peut être décomposé en une somme de sinusoïdes d'amplitudes, de fréquences et de phases différentes. Ce formalisme s'applique alors en considérant que le comportement du circuit est la superposition de sa réaction à chacune des sinusoïdes.

Aux bornes d'un dipôle D, l'impédance est définie comme le rapport $\underline{Z} = v(t)/i(t)$.



$$\text{avec } v(t) = V_0 \cdot \cos(\omega \cdot t + \varphi_V) \text{ et } i(t) = I_0 \cdot \cos(\omega \cdot t + \varphi_I)$$

On a donc $|\underline{Z}| = \frac{V_0}{I_0}$, $\text{Arg}(\underline{Z}) = \varphi = \varphi_V - \varphi_I$ et $\underline{Z} = \frac{V_0}{I_0} \cdot [\cos(\varphi) + j \cdot \sin(\varphi)]$

Dans les circuits linéaires, ce rapport ne dépend pas des caractéristiques de $v(t)$. L'accès d'un réseau est caractérisé par une paire de branches sur laquelle est appliqué un courant ou une tension. Les courants en A et en B sont identiques au signe près.

Attention

Il ne faut jamais utiliser la méthode de calcul complexe pour le calcul des puissances et des énergies car **cette notation n'est valable que pour des opérateurs linéaires** (addition, soustraction, intégrale, dérivée).

3. EN PRATIQUE

L'impédance et son inverse, l'admittance, sont des valeurs complexes qui sont soit exprimés en termes de modules et arguments, soit en écrivant leur partie réelle et leur partie imaginaire. Le tableau suivant présente les termes et les équivalences associés à chacun de ces éléments.

	Définitions	Termes	Équivalences
Impédance	$\underline{Z} = R + j \cdot X$	$R = \Re(\underline{Z})$: résistance	$R = \frac{G}{ \underline{Y} ^2}$
		$X = \Im(\underline{Z})$: réactance	$X = \frac{-B}{ \underline{Y} ^2}$
		$ \underline{Z} = \sqrt{R^2 + X^2}$	$\underline{Z} = \frac{1}{\underline{Y}} = \frac{1}{ \underline{Y} } \cdot e^{j(-\psi)}$
		$\varphi = \text{Arctan}\left(\frac{X}{R}\right)$	$\varphi = -\psi$
Admittance	$\underline{Y} = G + j \cdot B$	$G = \Re(\underline{Y})$: conductance	$G = \frac{R}{ \underline{Z} ^2}$
		$B = \Im(\underline{Y})$: susceptance	$B = \frac{-X}{ \underline{Z} ^2}$
		$ \underline{Y} = \sqrt{G^2 + B^2}$	$\underline{Y} = \frac{1}{\underline{Z}} = \frac{1}{ \underline{Z} } \cdot e^{j(-\phi)}$
		$\psi = \text{Arctan}\left(\frac{B}{G}\right)$	$\psi = -\phi$

8 Signaux périodiques

Mots clés

Courant alternatif, signaux sinusoïdaux, valeur moyenne, valeur efficace, période, fréquence, Hertz, phase.

1. EN QUELQUES MOTS

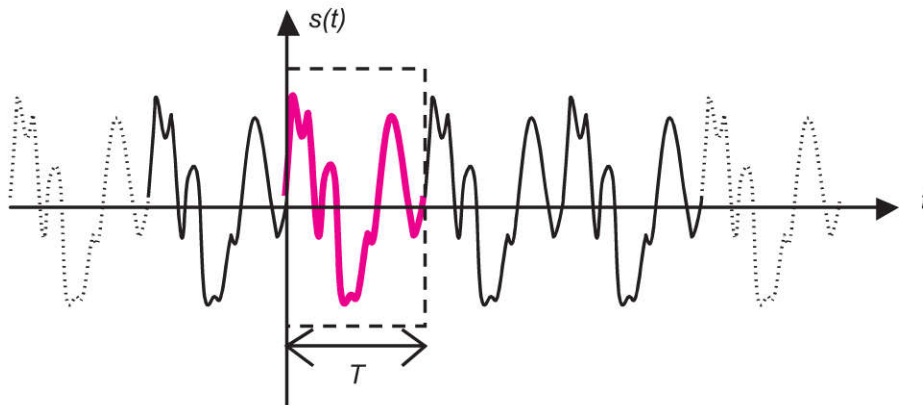
Les signaux périodiques sont incontournables dans l'analyse des circuits électroniques : utilisation des oscilloscopes (fiche 9), horloges, PLL (fiche 64), modulations (fiches 66 et 67) ou encore décomposition en séries de Fourier (fiche 15).

2. CE QU'IL FAUT RETENIR

Une fonction est dite périodique de **période** T si elle répond à la relation suivante :

$$s(t) = s(t + n \cdot T) \quad n \in \mathbb{Z}$$

Cela signifie qu'après une durée T , la fonction se reproduit égale à elle-même. Par conséquent, la description de la fonction sur une période suffit à la définir sur l'intervalle $] -\infty, +\infty [$. L'unité utilisée pour la période est la **seconde**.



La **fréquence** indique le nombre de fois qu'un signal se reproduit égal à lui-même par unité de temps. Par convention, l'unité de temps utilisée est la **seconde** et l'unité de fréquence est le

Hertz : $1 \text{ Hz} = 1 \text{ s}^{-1}$

Si un phénomène périodique dure 0,2 s, cela signifie que le phénomène se reproduit à l'identique 5 fois par seconde, sa fréquence est de 5 Hz. Nous pouvons déduire de cette observation que la période est l'inverse de la fréquence :

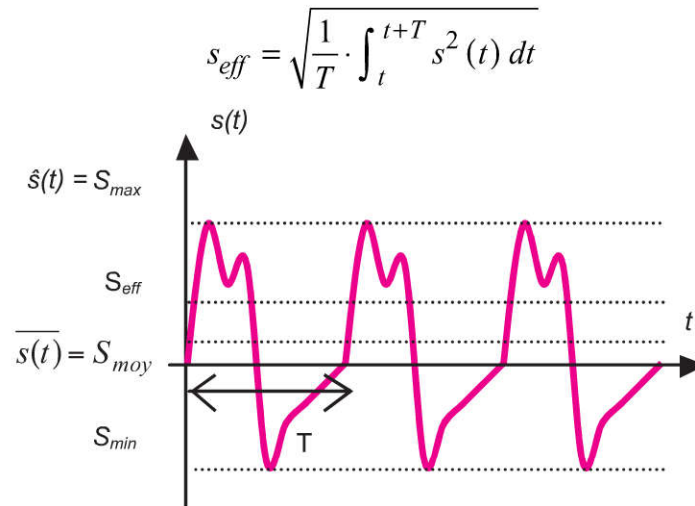
$$f = \frac{1}{T} \Leftrightarrow T = \frac{1}{f}$$

- La **valeur crête**, notée $\hat{s}(t)$, correspond au maximum qu'atteint la courbe. Pour une fonction périodique, $\hat{s}(t) = \hat{s}(t + n \cdot T)$.
- La **valeur crête-à-crête** correspond à la différence entre la valeur maximale et la valeur minimale du signal.

- La **valeur moyenne** d'un signal est également appelée **composante continue**. En effet, c'est une valeur constante : elle ne dépend pas du temps. Pour un signal périodique, elle s'exprime de la manière suivante :

$$\bar{s} = \frac{1}{T} \cdot \int_t^{t+T} s(t) dt$$

- La **valeur efficace**, s'appelle également **valeur RMS** (*Root Mean Square*) en référence à l'expression mathématique qui permet de la calculer. La valeur efficace du courant ou de la tension correspond à la valeur continue en courant ou en tension qui dissiperait la même puissance dans une résistance.



a) Signaux sinusoïdaux

Les tensions produites par les générateurs alternatifs fournissant l'énergie au réseau électrique domestique sont sinusoïdales, d'où leur importance en électricité et en électronique. De plus, les fonctions sinusoïdales sont utilisées pour mathématiser de nombreux phénomènes et formalismes électroniques qui sont traités dans ces fiches (électromagnétisme, télécommunications, séries de Fourier...).

$$s(t) = A \cdot \sin(\omega \cdot t + \phi_0)$$

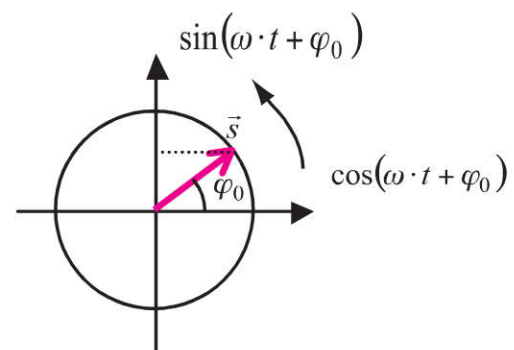
La valeur crête correspondant à l'amplitude A d'une sinusoïde, la valeur crête à crête vaut donc $2A$.

3. EN PRATIQUE

Le **cercle trigonométrique** est une représentation graphique qui permet d'illustrer les différents paramètres des fonctions sinusoïdales. Les valeurs instantanées de signaux sinusoïdaux correspondent à la projection du vecteur $\vec{s}$ sur les axes.

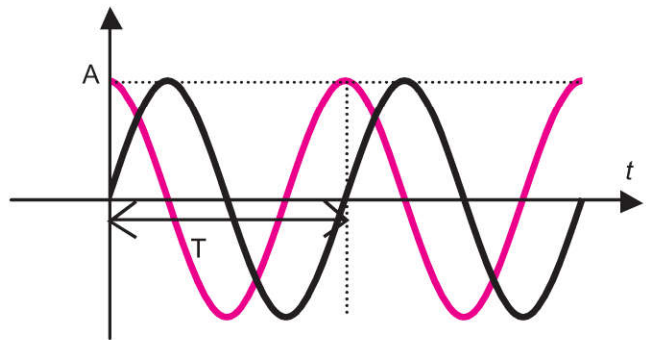
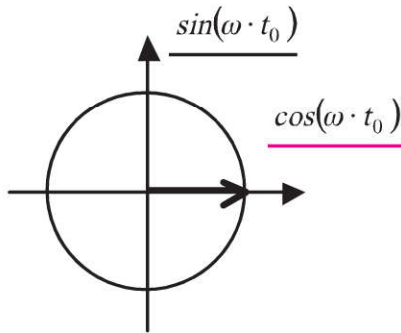
La **pulsation**, aussi appelée **vitesse angulaire**, indique la vitesse à laquelle le vecteur $\vec{s}$ parcourt le cercle trigonométrique. Cette grandeur s'exprime en rad.s^{-1} .

La **phase à l'origine** correspond à l'angle formé par le vecteur $\vec{s}$ et l'axe des abscisses à l'instant $t_0 = 0$.

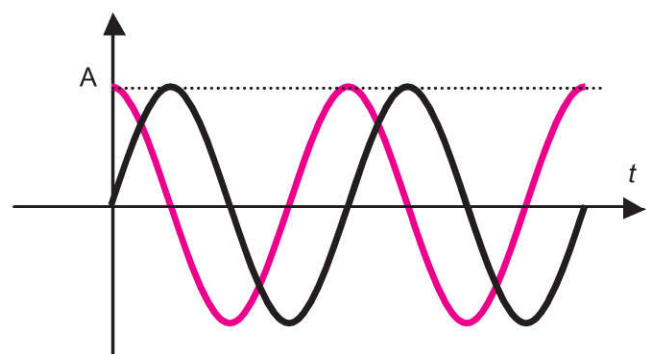
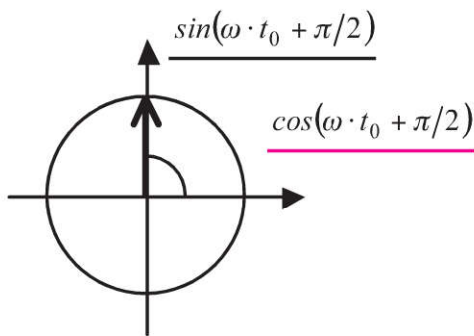


Exemple 1 : $\varphi_0 = 0$

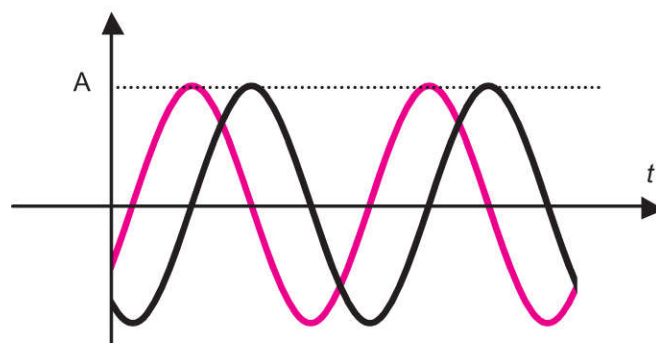
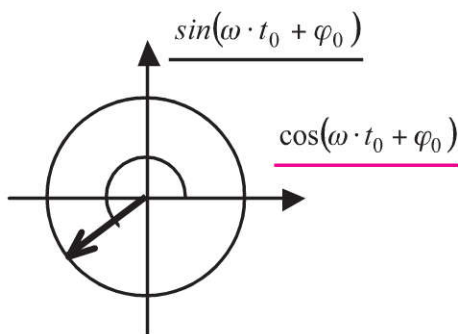
$$s(t) = A \cdot \cos\left(\omega \cdot t + \frac{\pi}{2}\right) = A \cdot \sin(\omega \cdot t)$$

**Exemple 2 :** $\varphi_0 = \frac{\pi}{2}$

$$s(t) = A \cdot \cos(\omega \cdot t) = A \cdot \sin\left(\omega \cdot t - \frac{\pi}{2}\right)$$

**Exemple 3 :** $\pi < \varphi_0 < \frac{3\pi}{2}$

$$s(t) = A \cdot \cos\left(\omega \cdot t + \varphi_0 + \frac{\pi}{2}\right) = A \cdot \sin(\omega \cdot t + \varphi_0)$$



La phase à l'origine peut être vue comme un décalage temporel t_0 (avance ou retard) appliqué à la fonction sinus.

$$s(t) = \sin(\omega \cdot t + \varphi_0) = \sin[\omega \cdot (t + t_0)] \Rightarrow t_0 = T \cdot \frac{\varphi}{2\pi} = \frac{\varphi}{\omega}$$

On parle de retard ou d'avance de phase lorsque deux sinusoides de même fréquence n'ont pas la même phase à l'origine. En se reportant à l'exemple 1 ci-dessus, on constate que la fonction $\cos(\omega \cdot t_0)$ est en **avance de phase** de $\pi/2$ sur la fonction $\sin(\omega \cdot t_0)$ et réciproquement, la fonction $\sin(\omega \cdot t_0)$ est en **retard de phase** de $\pi/2$ sur la fonction $\cos(\omega \cdot t_0)$.

Dans le cas particulier des signaux sinusoïdaux, la **valeur efficace** se calcule de la façon suivante :

$$s_{eff} = \sqrt{\frac{1}{T} \cdot \int_0^T s_{max}^2 \cdot \sin^2(2\pi \cdot f \cdot t) dt}$$

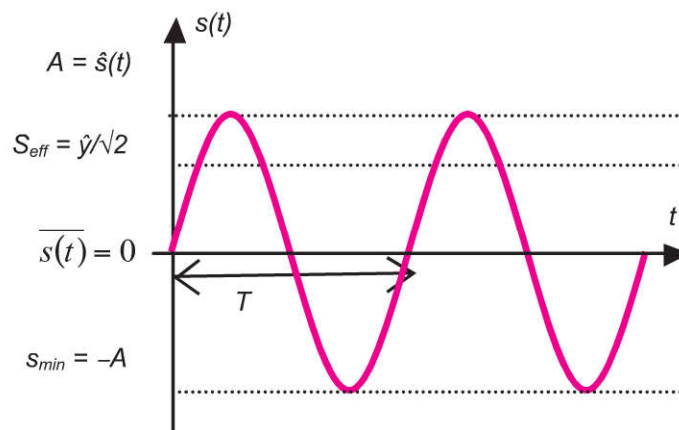
$$= \sqrt{\frac{1}{T} \cdot s_{max}^2 \cdot \int_0^T \sin^2(\theta) d\theta}$$

$$\text{Avec } \int \sin^2(\theta) d\theta = \frac{2 \cdot \theta - \sin(2 \cdot \theta)}{4} + k \Rightarrow \int_0^T \sin^2(2\pi \cdot f \cdot t) dt = \frac{T}{2}$$

$$\text{On obtient } s_{eff} = \frac{s_{max}}{\sqrt{2}}$$

En rapport avec la définition de la valeur efficace, on peut dire que courant domestique alternatif sinusoïdal de 230 V efficaces correspond à une valeur crête de 325 V :

$$y_{max} = y_{eff} \cdot \sqrt{2} \approx 230 \times 1,41 \approx 325 \text{ V}$$



Symbole	Nom	Unités	Signification	Équivalences
ω	Pulsation	rad.s^{-1}	Vitesse angulaire	$\omega = 2\pi \cdot f = \frac{2\pi}{T}$
f	Fréquence	Hz	Nombre de tours par seconde sur le cercle trigonométrique	$f = \frac{1}{T} = \frac{\omega}{2\pi}$
T	Période	s	Durée d'un tour sur le cercle trigonométrique	$T = \frac{1}{f} = \frac{2\pi}{\omega}$
φ_0	Phase à l'origine	Rad	Phase à l'instant $t = 0$	$t_0 = T \cdot \frac{\varphi_0}{2\pi} = \frac{\varphi_0}{\omega}$
φ	Phase instantanée	Rad	Phase à l'instant t	$\varphi = \omega \cdot t + \varphi_0 = 2\pi \cdot f \cdot t + \varphi_0$

Remarque

Sur la base d'une fonction sinusoïdale, l'amplitude, la fréquence et la phase peuvent varier dans le temps, on parle alors de modulation. Ce sujet est traité dans les [fiches 66](#) et [67](#).

9 Utilisation des oscilloscopes

Mots clés

Balayage horizontal, balayage vertical, calibre, fréquence, amplitude, déphasage.

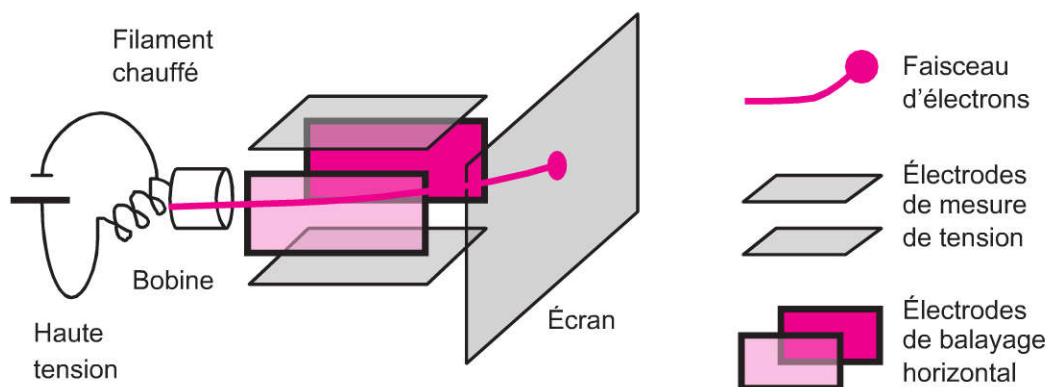
1. EN QUELQUES MOTS

Un oscilloscope, quelle que soit sa technologie (tube cathodique ou affichage LCD) ou sa marque, permet de visualiser et de mesurer des tensions. Cet appareil trouve principalement son intérêt dans l'analyse de signaux périodiques (fiche 8).

2. CE QU'IL FAUT RETENIR

Pour appréhender le fonctionnement des oscilloscopes, penchons-nous tout d'abord sur le **principe des tubes cathodiques**.

Le dispositif présenté ici est entièrement encapsulé dans une enceinte vidée de tout gaz. Un filament est chauffé grâce à une source de haute tension, ce qui a pour conséquence de libérer des électrons dans l'espace situé autour de lui. Ces électrons sont accélérés et focalisés à l'aide de cathodes et d'anodes créant un champ électrique. Le faisceau ainsi produit est orienté à l'aide de plaques parallèles soumises à une **différence de potentiel** : plus la différence de potentiel est élevée et plus le faisceau est fortement dévié. Les électrons frappent finalement une plaque de verre recouverte d'une couche de matériau fluorescent ce qui crée un point lumineux. Sur un oscilloscope, l'écran présente un quadrillage et des graduations.



La différence de potentiel appliquée sur les électrodes de mesure de tension est pilotée par le signal à observer. Un oscilloscope possède généralement deux ou quatre voies (ou entrées), c'est-à-dire qu'il permet d'afficher jusqu'à deux ou quatre signaux analogiques simultanément. La **sensibilité verticale** est sélectionnée par l'utilisateur et s'exprime en **Volts par division** (V/div.) ou en sous-multiple du Volt par division (mV/div., μ V/div.). Concernant la mesure des tensions, il existe les modes AC et DC. En mode DC, l'amplificateur vertical laisse passer la composante continue (fiche 8), alors que celle-ci est filtrée en mode AC.

La différence de potentiel appliquée sur les électrodes de balayage horizontal est périodique : elle fait circuler le faisceau d'électrons de gauche à droite plus ou moins rapidement, en fonction de la vitesse de balayage sélectionnée par l'utilisateur.

La **vitesse de balayage** correspond au temps qu'il faut à l'oscilloscope pour parcourir une graduation horizontale. L'unité de ce calibre est indiquée en **secondes par division** (s/div.) ou en sous-multiple de la seconde par division (ms/div., μ s/div., ns/div.).

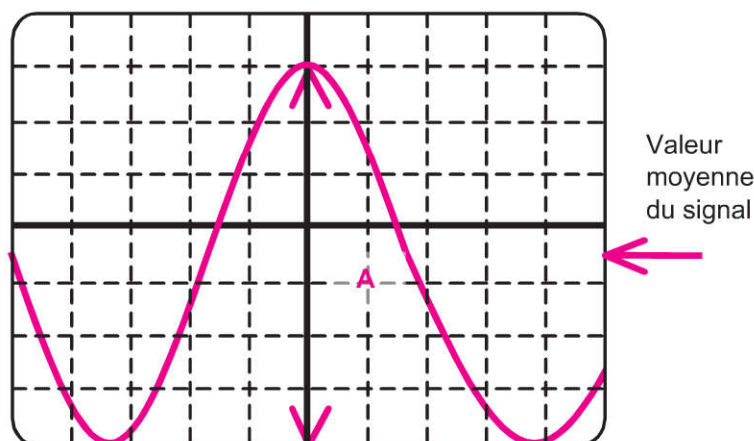
Remarque

Ces explications permettent d'aborder intuitivement le fonctionnement de l'oscilloscope et notamment la notion de calibre. Un grand nombre d'oscilloscopes actuels sont numériques et utilisent un affichage LCD, mais les principes de balayage horizontal et la sensibilité verticale restent valides.

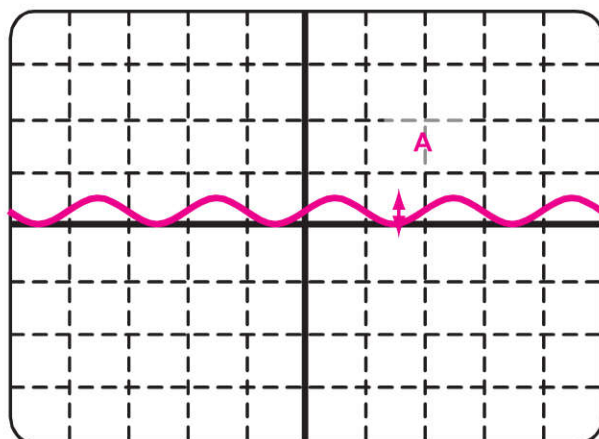
3. EN PRATIQUE

Un oscilloscope permet d'effectuer différentes mesures sur des signaux périodiques dont les plus courantes sont la mesure de l'amplitude, de la fréquence et du déphasage (fiche 8).

La **mesure de l'amplitude** d'un signal est optimale lorsque celui-ci est étiré au maximum sur l'axe des ordonnées. La technique illustrée ici consiste à décaler le signal sur l'axe vertical de sorte que le minimum soit tangent à la graduation la plus basse. Il suffit alors de centrer la valeur maximum sur l'axe vertical. La lecture du nombre de graduations indique l'amplitude relative du signal (ici : 7 carreaux environ). Pour obtenir la valeur en volts, il faut utiliser la sensibilité verticale sélectionnée. Par exemple, si la sensibilité S est de 200 mV/div., la valeur d'amplitude mesurée est de 1,4 V.



L'oscillogramme suivant illustre la difficulté de mesurer l'amplitude sur un signal mal calibré : l'erreur de mesure peut être de quelques dizaines de pour cents.

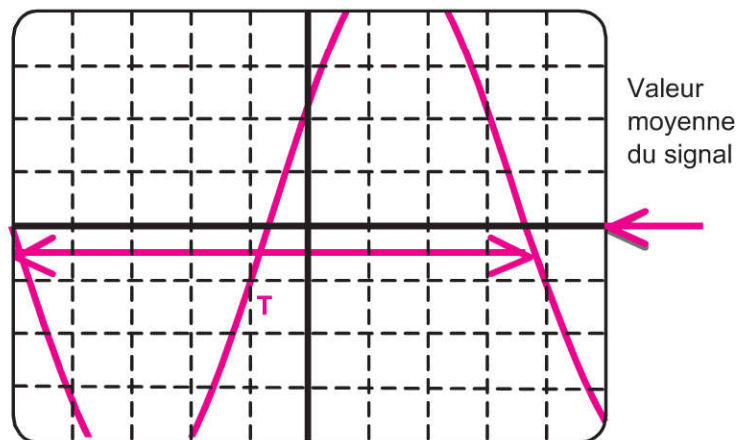


Rappel

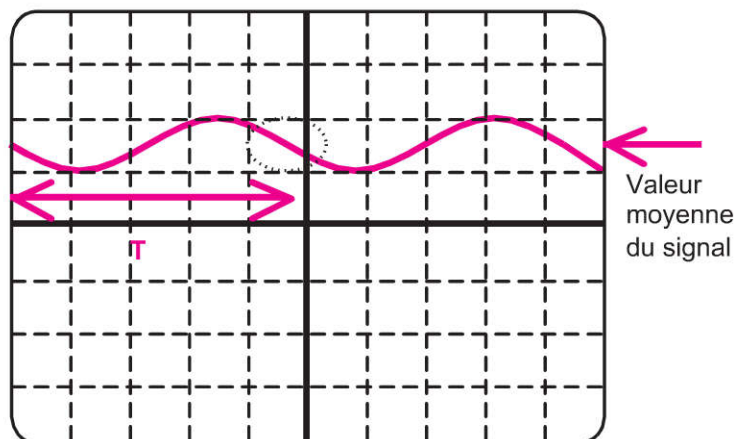
Les oscilloscopes servent à visualiser des signaux périodiques. Afin de **mesurer la fréquence** d'un signal, il convient de **choisir les calibres tels qu'une seule période de signal complète soit affichée** sur l'axe horizontal. Le signal doit être centré et dilaté verticalement afin d'obtenir des pentes raides aux passages par zéro. La

distance séparant deux passages par zéro et définissant une période est convertie en durée. L'inverse de cette durée est la période. Dans notre exemple, admettons qu'un carreau corresponde à $2\text{ }\mu\text{s}$. La période T occupant

environ 8,6 divisions vaut alors $17,2\text{ }\mu\text{s}$. La fréquence est donc de $f = \frac{1}{T} \approx \frac{1}{17,2 \cdot 10^{-6}} \approx 58\text{ kHz}$



Dans l'exemple suivant, le signal n'est ni centré ni dilaté, il est donc difficile d'évaluer le passage par zéro à l'intérieur de la zone délimitée par les pointillés.



Mesurer un déphasage n'a de sens qu'en considérant **deux signaux de même fréquence**. La technique est similaire à la mesure de fréquence : il faut que les passages par zéro des deux signaux soient nets et que les signaux soient correctement centrés. La distance séparant ces deux signaux est proportionnelle au déphasage.

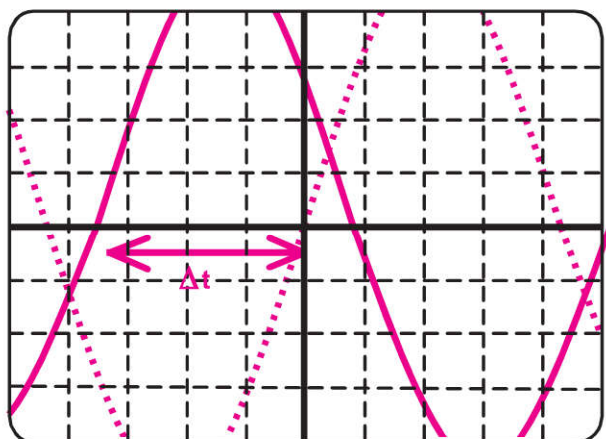
Dans l'exemple proposé, 3,5 carreaux séparent temporellement les deux courbes. En reprenant le calibre de $2\text{ }\mu\text{s/div.}$ du cas précédent, la courbe en pointillés présente un retard Δt de $7\text{ }\mu\text{s}$ par rapport à la courbe en trait plein. La période étant toujours de $17,2\text{ }\mu\text{s}$, le déphasage relatif en radians $\Delta\phi$ se calcule comme suit :

$$\Delta\phi = \frac{\Delta t}{T} \cdot 2 \cdot \pi \approx \frac{7 \cdot 10^{-6}}{17,2 \cdot 10^{-6}} \cdot 2 \cdot \pi \approx 2,56 \text{ rad (mod } 2 \cdot \pi)$$

Il suffit de remplacer 2π par 360 pour obtenir ce déphasage en degrés.

Remarque

Il est impossible de savoir combien de périodes séparent les deux signaux, le résultat est donc indiqué à $2 \cdot k \cdot \pi$ près avec $k \in \mathbb{Z}$.



Les oscilloscopes les plus modernes sont numériques et proposent d'effectuer automatiquement la plupart des mesures courantes (amplitude, fréquence, déphasage...) et de calibrer automatiquement la courbe visualisée. Certains peuvent être considérés comme de petits ordinateurs dédiés à la mesure et dotés d'un système d'exploitation. Cette famille d'oscilloscopes propose d'effectuer des mesures, d'enregistrer des courbes et de gérer les fichiers correspondants de façon intuitive.

Leurs performances peuvent être évaluées selon différents paramètres dont les principaux sont :

- le nombre de voies de mesure ;
- la fréquence maximale d'entrée mesurable. Pour les oscilloscopes numériques, il faut également considérer la fréquence d'échantillonnage maximale (en échantillons par seconde, [fiche 61](#)) ;
- la possibilité de sauvegarder les mesures, sur un disque dur interne, une clé USB ou un ordinateur relié à l'oscilloscope.

10 Lois de Kirchhoff

Mots clés

Réseau, branche, loi de nœuds, loi des mailles.

1. EN QUELQUES MOTS

Les lois de Kirchhoff, également appelées **loi des nœuds** et **loi des mailles**, sont à la base de l'analyse de tout circuit électrique et électronique par extension. Elles permettent de calculer la répartition des courants et des tensions dans un réseau, ce qui est fondamental dans le domaine étudié dans cet ouvrage.

2. CE QU'IL FAUT RETENIR

a) Définitions

Un **réseau** est un ensemble de générateurs et de récepteurs qui constituent des circuits.

Un réseau est dit **linéaire** lorsque les lois qui le régissent sont linéaires.

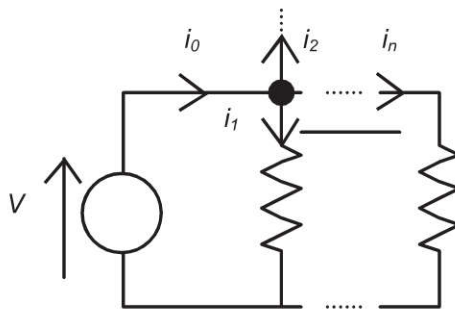
Une **branche** est un ensemble d'éléments montés en série.

Un **nœud** est un point commun à trois branches au moins.

b) Loi de nœuds

La somme des courants qui entrent par un nœud est égale à la somme des courants qui en repartent. Cette loi est imposée par la conservation de la charge électrique : les charges ne peuvent pas s'accumuler sur un nœud qui est la convergence de plusieurs fils. Les charges qui arrivent à un nœud sont au même nombre que celles qui en repartent.

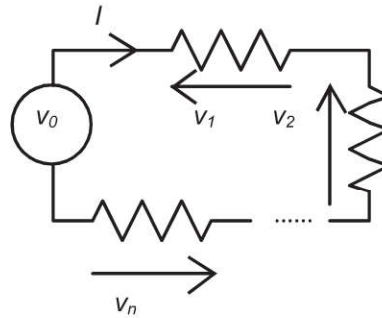
$$\sum_n i_n(t) = 0$$



c) Loi des mailles

La somme des tensions le long d'une maille quelconque d'un réseau est nulle.

$$\sum_n v_n(t) = 0$$

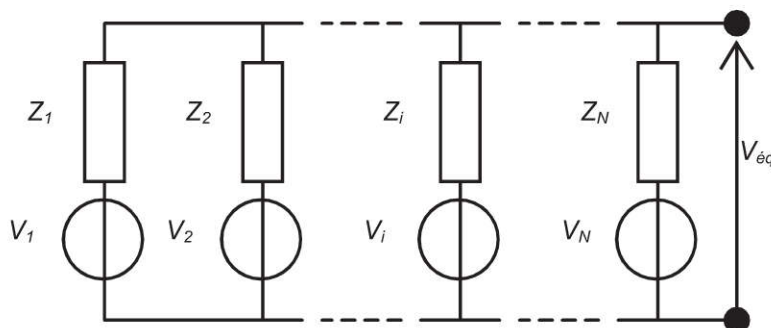


Afin d'appliquer la loi des mailles, il convient de procéder par étapes.

- Choisir un sens pour les courants dans les branches constituant la maille considérée.
- Choisir arbitrairement un sens de parcours.
- Il faut alors compter positivement les tensions aux bornes des résistances lorsque les courants sont dans le sens des courants choisis et négativement dans le cas contraire. Cette instruction découle des indications données au b.
- Compter positivement les forces électromotrices ou les forces contre électromotrices lorsqu'on rentre par le + et négativement lorsqu'on rentre par le –.
- Si l'on trouve un courant négatif, changer le signe du courant ainsi que celui du ou des récepteur(s) dans la branche.

d) Théorème de Millman

Le théorème de Millman représente un cas particulier de la loi des nœuds. Il donne la tension équivalente d'un réseau électrique de N branches parallèles constituées chacune d'un générateur de tension idéal et d'une impédance en série.



$$V_{eq} = \frac{\sum_{k=1}^N Y_k \cdot V_k}{\sum_{k=1}^N Y_k} = \frac{\sum_{k=1}^N \frac{V_k}{Z_k}}{\sum_{k=1}^N \frac{1}{Z_k}} \text{ avec } Y_k = \frac{1}{Z_k}$$

e) Solvabilité

Les lois de Kirchhoff permettent d'obtenir plusieurs équations pour un même réseau.

Si le réseau considéré comporte N nœuds et B branches, il existe B courants différents dans le réseau. Il faut donc écrire B équations indépendantes du réseau afin de calculer chaque courant ce qui revient à devoir constituer un système de B équations à B inconnues soit :

- $(N - 1)$ équations de nœuds ;
- $(B - N + 1)$ équations de maille.

11 Loi d'Ohm

Mots clés

Résistance, courant alternatif, régime sinusoïdal, notation complexe.

1. EN QUELQUES MOTS

La loi d'Ohm est une relation linéaire entre la tension aux bornes d'un dipôle et le courant le traversant.

Dans sa forme la plus élémentaire, elle concerne la relation tension-courant en continu ou en régime alternatif d'un dipôle résistif.

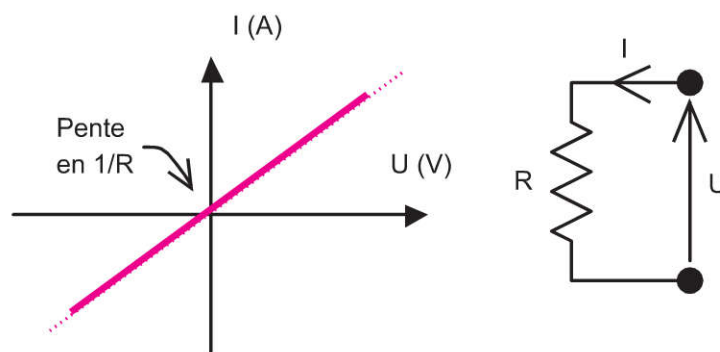
Dans sa version complexe, elle lie une tension sinusoïdale à un courant sinusoïdal de tout dipôle passif.

2. CE QU'IL FAUT RETENIR

Dans son expression la plus simple, la loi d'Ohm relie la tension et le courant aux bornes d'une résistance (la résistance s'exprime en Ohms, noté Ω). En voici la formulation **en courant continu** :

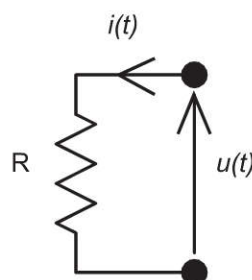
$$U = R \cdot I \Leftrightarrow I = U/R \Leftrightarrow R = U/I$$

Cela signifie que le courant aux bornes d'une résistance est proportionnel à la tension (en valeur absolue). On dit alors que **la caractéristique d'une résistance est linéaire** : c'est une droite de pente $1/R$ et d'ordonnée à l'origine nulle.

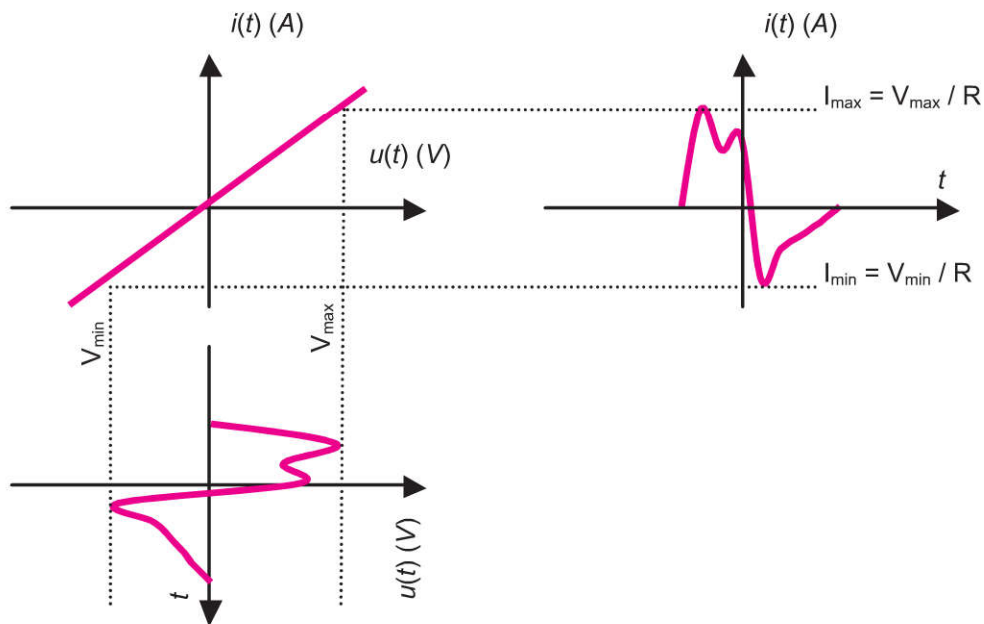


La loi d'Ohm est également applicable aux courants alternatifs :

$$u(t) = R \cdot i(t) \Leftrightarrow i(t) = u(t)/R$$

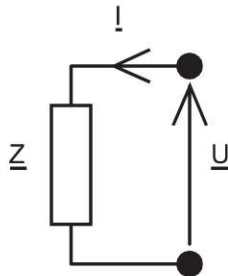


Théoriquement, la caractéristique d'une résistance ne change pas en régime alternatif, nous pouvons donc reprendre la figure précédente pour représenter le courant alternatif traversant une résistance et la tension à ses bornes.



La formulation la plus générale de la loi d'Ohm utilise la notation complexe et s'applique non seulement aux résistances mais aussi à tout dipôle passif d'impédance $\underline{Z}$.

$$\underline{U} = \underline{Z} \cdot \underline{I} \text{ avec } \underline{Z} = R + j \cdot X$$



12 Énergie, puissance

Mots clés

Joule, Watt, VA, puissance moyenne, puissance active, puissance réactive, théorème de Boucherot, facteur de puissance, facteur de qualité.

1. EN QUELQUES MOTS

Comme pour tous les domaines de la physique, les notions d'énergie et de puissance sont essentielles : l'énergie se manifeste par la production de mouvement, de lumière (onde électromagnétique), de chaleur. En électronique et en électricité, l'énergie électrique provenant d'un générateur ou l'énergie électromagnétique captée par une antenne provoque, ou induit, un déplacement d'électrons. Cette énergie est véhiculée à travers le circuit jusqu'aux charges dans lesquelles elle sera utilisée.

2. CE QU'IL FAUT RETENIR

a) Unités

L'unité internationale utilisée pour l'**énergie** est le **Joule**. Elle est équivalente à l'énergie déployée pour appliquer une force d'un Newton sur un mètre.

$$J = N \cdot m = \text{kg} \cdot \text{m}^2 \cdot \text{s}^{-2}$$

La **puissance** représente la quantité d'énergie dissipée ou reçue par un système par unité de temps. L'unité internationale est le **Watt** qui correspond à une énergie d'un Joule dissipé (ou reçu) en une seconde.

$$W = J \cdot \text{s}^{-1}$$

b) Calculs d'énergie et de puissance

La puissance est égale à la dérivée de l'énergie par rapport au temps (et réciproquement, l'énergie est l'intégrale de la puissance). Par conséquent, la **valeur instantanée de l'énergie dissipée** par la puissance $p(t)$ pendant la durée Δt entre les instants t_0 et t_1 peut s'écrire :

$$E = \int_{t_0}^{t_1} p(t) dt = \int_{t_0}^{t_1} u(t) \cdot i(t) dt$$

où $u(t)$ et $i(t)$ sont la tension et le courant instantanés.

Nous pouvons en déduire l'expression de la **puissance moyenne** dissipée ou reçue pendant

l'intervalle de temps Δt : $\bar{P} = \frac{E}{\Delta t} = \frac{1}{\Delta t} \int_{t_0}^{t_1} p(t) dt$ avec $\Delta t = t_1 - t_0$

En **régime périodique**, dans le cas général, il suffit d'effectuer le calcul précédent sur une seule période de signal pour obtenir la puissance moyenne. Δt devient donc T et t_1 devient $t_0 + T$:

$$\bar{P} = \frac{E}{T} = \frac{1}{T} \int_{t_0}^{t_0+T} p(t) dt$$

Dans le cadre du régime périodique, étudions le cas particulier du **régime permanent sinusoïdal**. Soient $v(t)$ la tension entre deux points d'un circuit non linéaire et $i(t)$ le courant circulant dans la branche considérée. En toute généralité, le caractère non linéaire du circuit induit un déphasage entre le courant et la tension.

$$\begin{cases} v(t) = \hat{V} \cdot \cos(\omega \cdot t) = \sqrt{2} \cdot V_{eff} \cdot \cos(\omega \cdot t) \\ i(t) = \hat{I} \cdot \cos(\omega \cdot t + \phi) = \sqrt{2} \cdot I_{eff} \cdot \cos(\omega \cdot t + \phi) \end{cases}$$

$$p(t) = 2 \cdot V_{eff} \cdot I_{eff} \cdot \cos(\omega \cdot t) \cdot \cos(\omega \cdot t + \phi) = V_{eff} \cdot I_{eff} \cdot [\cos(2 \cdot \omega \cdot t + \phi) + \cos(\phi)]$$

$$\overline{P} = \frac{1}{T} \int_0^T p(t) dt = \frac{V_{eff} \cdot I_{eff}}{T} \cdot \left\{ \frac{1}{2 \cdot \omega} [\sin(2 \cdot \omega \cdot t + \phi)]_0^T + \cos(\phi) \cdot T \right\} = V_{eff} \cdot I_{eff} \cdot \cos(\phi)$$

Le **théorème de Boucherot** permet de calculer la puissance totale consommée par un circuit électrique comportant plusieurs dipôles de facteurs de puissance différents et d'en déduire le courant total y circulant sans avoir recours à un diagramme de Fresnel. Seules les résistances dissipent de la puissance, et quel que soit le circuit, on additionne les puissances dissipées. Auparavant, il faut calculer pour cela la puissance apparente et la puissance réactive afin d'en déduire la puissance active, seule grandeur ayant un sens physique et à laquelle est attribuée l'unité Watt.

La puissance réactive est la composante imaginaire pure de la puissance dissipée aux bornes d'un dipôle dont l'impédance complexe s'écrit $\underline{Z} = R + j \cdot X$:

$$P_R = V_{eff} \cdot I_{eff} \cdot \sin(\phi) = |X| \cdot I_{eff}^2$$

La puissance réactive s'exprime en VA réactifs.

La **puissance apparente**, associée à l'unité VA (Volts-Ampères), est le simple produit des tensions et courants efficaces : $P_A = V_{eff} \cdot I_{eff}$

Nous pouvons alors déduire de ces calculs la valeur de la **puissance active** P

$$P_A^2 = P^2 + P_R^2 \Rightarrow P = V_{eff} \cdot I_{eff} \cdot \cos(\phi) = R \cdot I_{eff}^2$$

Remarque

$$\sin(\phi) = \frac{P_R}{P_A}, \cos(\phi) = \frac{P}{P_A} \text{ et donc } \tan(\phi) = \frac{P_R}{P}$$

- La **puissance complexe** est la somme de la puissance active (nombre réel) et de la puissance réactive (nombre imaginaire pur) : $\underline{P} = P + j \cdot P_R$
- Le **facteur de puissance** se définit comme étant le rapport entre la puissance active et la puissance apparente, on a donc : $f = \frac{P}{P_A} = \cos(\phi)$
- Le **facteur de qualité**, notion également abordée dans les [fiches 26](#) et [27](#) sur le filtrage, peut se définir comme le rapport du module de la puissance réactive et de la puissance active :

$$Q = \frac{|P_R|}{P_A}$$

3. EN PRATIQUE

La résistance absorbe de la puissance et ne la restitue jamais. Elle transforme de l'énergie de façon irréversible, c'est un élément dissipatif.

Dans le **cas particulier** du calcul de puissance pour **une inductance ou une capacité pure** (c'est-à-dire sans élément résistif), le déphasage noté ϕ vaut :

- $\pi/2$ dans le cas d'une capacité ;
- $-\pi/2$ dans le cas d'une inductance.

Dans les deux cas, $\cos(\phi)$ vaut 0, **la puissance dissipée est donc nulle**. La capacité et l'inductance peuvent restituer à tout moment l'énergie qu'elles ont absorbé, ce sont des éléments non dissipatifs qui emmagasinent de l'énergie.

13 Notation et transformation complexes

Mots clés

Fonction exponentielle, diagramme de Fresnel.

1. EN QUELQUES MOTS

La **notation complexe** permet à l'électronicien d'appliquer facilement à des grandeurs correspondant à des fonctions sinusoïdales des opérations mathématiques courantes :

- les quatre opérations élémentaires (+ ; - ; ∞ ; $\div$) ;
- dérivation et intégration.

2. CE QU'IL FAUT RETENIR

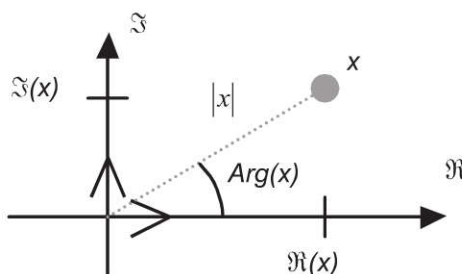
Soit x un nombre complexe. Rappelons qu'un nombre complexe peut être déterminé de deux façons :

1. soit en donnant sa partie réelle et sa partie imaginaire : $x = \Re(x) + j \cdot \Im(x)$;
2. soit en donnant son module et son argument et en utilisant la notation exponentielle :

$$x = |x| \cdot e^{j \cdot \text{Arg}(x)}.$$

Rappelons la relation suivante : $e^{j \cdot \theta} = \cos(\theta) + j \cdot \sin(\theta)$

La représentation graphique reprend les notations utilisées dans les équations de cette fiche.



Nous constatons que cette dernière notation contient les données caractérisant une fonction sinusoïdale :

- l'amplitude $\rightarrow$ le module
- l'argument $\rightarrow$ la phase instantanée

Le formalisme des nombres complexes est très efficace pour effectuer des calculs sur les signaux sinusoïdaux car un seul nombre suffit à représenter le module et la phase.

Amplitude	A
Fréquence	F
Pulsation	ω
Phase à l'origine	φ
Phase instantanée	$\omega \cdot t + \varphi = 2 \cdot \pi \cdot f \cdot t + \varphi$

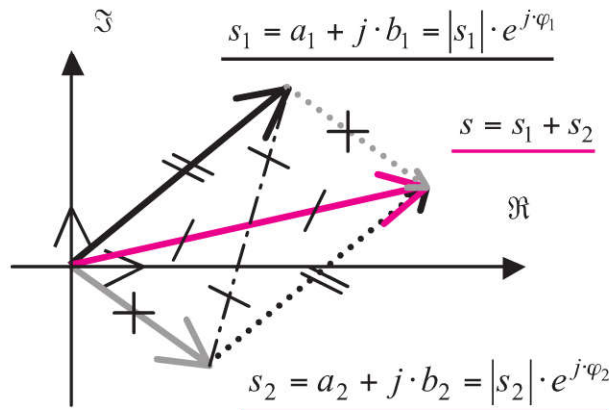
Signal sinusoïdal	$S(\omega) = A \cdot e^{j \cdot \omega \cdot t + \varphi}$
	$S(f) = A \cdot e^{j \cdot 2 \cdot \pi \cdot f \cdot t + \varphi}$
Dérivation	$S'(\omega) = A \cdot j \cdot \omega \cdot e^{j \cdot \omega \cdot t + \varphi} = j \cdot \omega \cdot S(\omega)$
Intégration	$\int S(\omega) d\omega = \frac{A}{j \cdot \omega} \cdot e^{j \cdot \omega \cdot t + \varphi} = \frac{1}{j \cdot \omega} \cdot S(\omega)$

3. EN PRATIQUE

Soit $s_1(t)$ et $s_2(t)$ deux sinusoïdes d'amplitudes et de phases différentes, mais de même fréquence (ou période). Ces fonctions ont pour équation :

$$\begin{cases} s_1(t) = |s_1| \cdot \cos(\omega \cdot t + \varphi_1) \\ s_2(t) = |s_2| \cdot \cos(\omega \cdot t + \varphi_2) \end{cases}$$

Connaissant leur fréquence (ou leur pulsation), nous pouvons représenter leur amplitude et leur phase respectives dans le **plan complexe**. C'est ce que l'on appelle le **diagramme de Fresnel**.



Il est alors aisé de représenter la somme des deux signaux dans ce plan sous forme de la somme

de deux vecteurs avec $|s| = \sqrt{|s_1|^2 + |s_2|^2}$ et $\text{Arg}(s) = \arctan\left(\frac{a_1 + a_2}{b_1 + b_2}\right)$

Attention

La multiplication de deux signaux sinusoïdaux de même fréquence ne peut faire appel au même type de construction : le résultat du produit des deux nombres complexes fait appel au produit des modules et à la somme des arguments :

$$s_1 \cdot s_2 = |s_1| \cdot |s_2| \cdot e^{j \cdot (\varphi_1 + \varphi_2)}$$

14 Fonctions utiles en traitement du signal

Mots clés

Unité, système international, ampère, volt, coulomb, watt, joule.

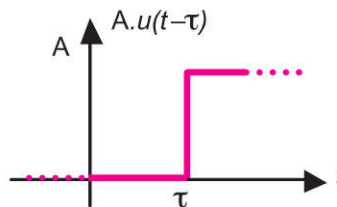
1. EN QUELQUES MOTS

Quelques fonctions élémentaires très employées en traitement du signal nous seront utiles par la suite : le saut unité, la fonction rectangle, l'impulsion de Dirac, le peigne de Dirac.

2. CE QU'IL FAUT RETENIR

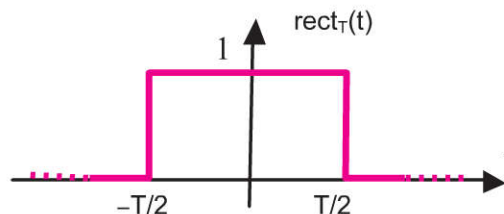
La **fonction échelon** ou saut unité notée $u(t)$ vaut 0 lorsque t est négatif et 1 lorsque t est positif.

$$u(t - \tau) = \begin{cases} 1, & t \geq \tau \\ 0, & t < \tau \end{cases}$$



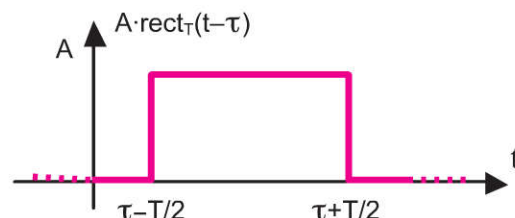
La **fonction rectangle**, ou **fonction porte**, notée $rect_T(t)$, vaut 1 entre $T/2$ et $T/2$ et 0 ailleurs. On peut également écrire :

$$rect_T(t) = \begin{cases} 1, & |t| < T/2 \\ 0, & |t| \geq T/2 \end{cases}$$

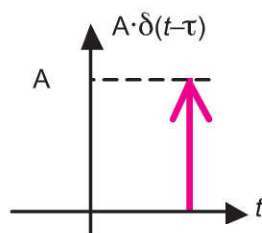


Cette fonction est paire, de largeur T , et son aire vaut T . La fonction rectangle peut être décomposée en deux fonctions saut unité : $rect_T(t) = u(t + T/2) \times u(-t - T/2)$

$$rect_T(t - \tau) = \begin{cases} 0, & t \leq \tau - T/2 \\ 1, & \tau - T/2 < t < \tau + T/2 \\ 0, & t \geq \tau + T/2 \end{cases}$$



L'**impulsion de Dirac** est un objet mathématique qui n'est pas à proprement parler une fonction mais une distribution. Une impulsion idéale est de **largeur nulle** et de **surface unité**. Par conséquent, *son amplitude est infinie*. C'est pour cette raison qu'on la représente sous forme d'une flèche.



Une impulsion de Dirac est en quelque sorte une fonction rectangle dont l'aire est de 1, dont la largeur tend vers 0 et dont l'amplitude tend vers l'infini pour conserver l'aire unité.

$$\delta(t - \tau) = T \cdot \text{rect}_{1/T}(t - \tau) \text{ lorsque } T \rightarrow \infty$$

Propriétés

Multiplier une fonction par un Dirac consiste à prélever un **échantillon** de cette fonction :

$$x(t) \cdot \delta(t) = x(0) \cdot \delta(t) \text{ et plus généralement, } x(t) \cdot \delta(t - \tau) = x(\tau) \cdot \delta(t - \tau)$$

L'impulsion de Dirac est la **dérivée de la fonction échelon**.

$$\frac{du(t)}{dt} = \delta(t) \text{ et plus généralement, } \frac{du(t - \tau)}{dt} = \delta(t - \tau)$$

On peut en déduire que

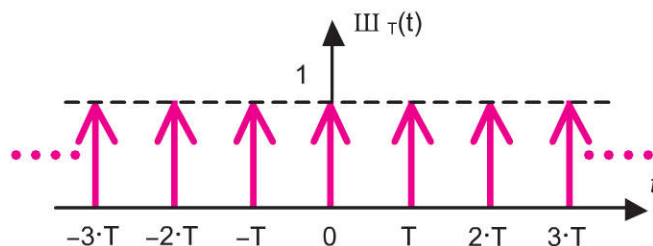
$$\int_{-\infty}^{+\infty} x(t) \cdot \delta(t - \tau) dt = x(\tau) \text{ car } \int_{-\infty}^{+\infty} x(t) \cdot \delta(t - \tau) dt = x(\tau) \cdot \int_{-\infty}^{+\infty} \delta(t - \tau) dt \text{ et}$$

$$\int_{-\infty}^{+\infty} \delta(t) dt = 1$$

La **convolution d'une fonction par un Dirac transpose** celle-ci sur l'axe des abscisses (fiche 18) :

$$x(t) * \delta(t - \tau) = x(t - \tau)$$

Un **peigne de Dirac**, noté $\text{III}_T(t)$ en référence à la forme de la lettre cyrillique « cha », est une suite périodique infinie de Dirac séparés par une durée T .



$$\text{III}_T(t) = \sum_{k=-\infty}^{+\infty} \delta(t - kT)$$

Cette fonction est particulièrement utile pour l'étude de l'échantillonnage dans les domaines temporel et fréquentiel (fiche 61).

Propriété

La transformée de Fourier d'un peigne de Dirac de période T est un peigne de Dirac de période $1/T$ et de poids $1/T$. (fiche 16)

$$\text{III}_T(t) \xrightarrow{\text{TF}} \frac{1}{T} \cdot \text{III}_{\frac{1}{T}}(f) \text{ soit } \sum_{k=-\infty}^{+\infty} \delta(t - k \cdot T) \xrightarrow{\text{TF}} \frac{1}{T} \cdot \sum_{k=-\infty}^{+\infty} \delta\left(f - \frac{k}{T}\right)$$

15 Séries de Fourier

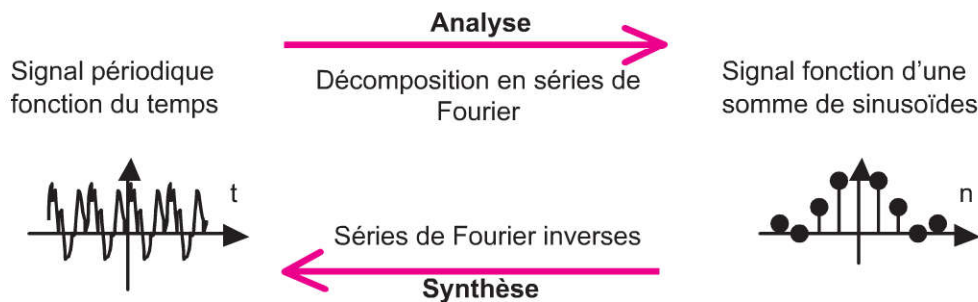
Mots clés

Signal périodique, analyse, synthèse, domaine fréquentiel, harmonique, fondamentale.

1. EN QUELQUES MOTS

Le développement en séries de Fourier est un outil mathématique qui permet de décomposer un **signal périodique** en une somme de sinusoïdes. Cela revient à considérer qu'un signal périodique est la somme de plusieurs sinusoïdes (voire une infinité) d'amplitudes, de phases et de fréquences différentes. Ces différentes phases et amplitudes constituent les **coefficients** de la série de Fourier.

Les différentes fréquences des sinusoïdes sont des **harmoniques** de la fréquence du signal périodique d'origine, aussi appelée fréquence **fondamentale** f_0 . Les fréquences harmoniques sont des multiples entiers de la fréquence fondamentale. Le coefficient multiplicateur, noté n , est le **rang de l'harmonique** considérée.



2. CE QU'IL FAUT RETENIR

Soit $s(t)$ un signal périodique quelconque. Sa période est notée T et sa fréquence f_0 . Celui-ci peut-être reconstitué en utilisant les coefficients calculés par série de Fourier. Cette opération de reconstitution du signal s'appelle la **synthèse**. Chaque coefficient est un nombre réel correspondant à l'amplitude d'une fonction sinusoïdale de fréquence $n \cdot f_0$.

$$s(t) = a_0 + \sum_{n=1}^{+\infty} [|c_n| \cdot \cos(n \cdot 2 \cdot \pi \cdot f_0 \cdot t + \varphi_n)] = \sum_{n=-\infty}^{+\infty} [|c_n| \cdot e^{j \cdot n \cdot 2 \cdot \pi \cdot f_0 \cdot t}]$$

$$= a_0 + \sum_{n=1}^{+\infty} [a_n \cdot \cos(n \cdot 2 \cdot \pi \cdot f_0 \cdot t) + b_n \cdot \sin(n \cdot 2 \cdot \pi \cdot f_0 \cdot t)]$$

Les coefficients a_n pondèrent des fonctions cosinus de fréquence $n \cdot f_0$ et les coefficients b_n pondèrent des fonctions sinus de fréquence $n \cdot f_0$. a_0 représente la composante continue du signal périodique, c'est-à-dire sa valeur moyenne : ce nombre est donc indépendant du temps.

La synthèse d'un signal correspond à une somme infinie, c'est cette somme infinie qui correspond au concept de série.

Le calcul des coefficients est l'opération d'**analyse** que nous allons détailler ici.

Les a_n sont calculés grâce à la formule suivante :

$$a_n = \frac{2}{T} \int_{-T/2}^{T/2} s(t) \cdot \cos\left(n \cdot 2 \cdot \pi \cdot \frac{t}{T}\right) dt$$

T est la période du signal $s(t)$.

Dans le cas particulier où $n = 0$, a_n devient

$$a_0 = \frac{2}{T} \int_{-T/2}^{T/2} s(t) \cdot \cos(0) dt = \frac{2}{T} \int_{-T/2}^{T/2} s(t) dt$$

a_0 étant la valeur moyenne du signal, cette expression est également présentée dans la [fiche 8](#).

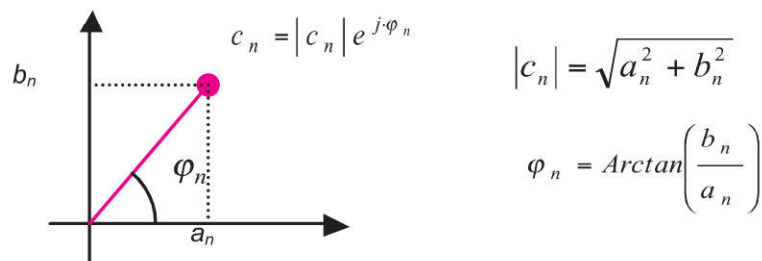
$$a_0 = \frac{2}{T} \int_{-T/2}^{T/2} s(t) dt = \overline{s(t)}$$

Dans le cas particulier où $n = 0$, b_n devient

$$b_n = \frac{2}{T} \int_{-T/2}^{T/2} s(t) \cdot \sin(0) dt = 0$$

$b_0 = 0$

Les coefficients a_n et b_n peuvent être combinés sous la forme d'un seul et unique nombre complexe c_n par simple addition : a_n est la partie réelle de c_n et b_n sa partie imaginaire pure.



Les coefficients complexes c_n peuvent, au même titre que les a_n et b_n , être calculés directement à partir de la fonction $s(t)$:

$$c_n = \frac{1}{T} \int_{-T/2}^{T/2} s(t) \cdot e^{-j \cdot 2 \cdot \pi \cdot \frac{n}{T} \cdot t} dt$$

Rappel : $e^{j \cdot x} = \cos(x) + j \cdot \sin(x) \Rightarrow e^{-j \cdot x} = \cos(x) - j \cdot \sin(x)$

Cas particuliers : fonctions paires et impaires

La décomposition en série de Fourier d'une fonction paire implique que les coefficients b_n sont nuls. En effet, les coefficients b_n sont ceux pondérant les fonctions *sinus* qui elles, sont impaires : elles ne contribuent donc pas dans la constitution du signal.

De même, la décomposition en série de Fourier d'une fonction impaire implique que les coefficients a_n sont nuls puisque les coefficients a_n pondèrent les fonctions *cosinus* qui sont paires.

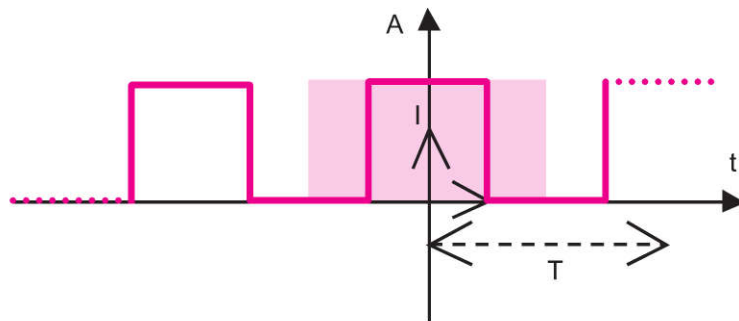
Coef.	Calcul	Interprétation	Cas particuliers
a_0	$= \frac{2}{T} \int_{-T/2}^{T/2} s(t) dt = \overline{s(t)}$	Valeur moyenne du signal	= 0 pour les fonctions sans composante continue
a_n	$= \frac{2}{T} \int_{-T/2}^{T/2} s(t) \cdot \cos\left(n \cdot 2 \cdot \pi \cdot \frac{t}{T}\right) dt$	$= \Re(c_n)$	= 0 pour les fonctions impaires
b_n	$= \frac{2}{T} \int_{-T/2}^{T/2} s(t) \cdot \sin\left(n \cdot 2 \cdot \pi \cdot \frac{t}{T}\right) dt$	$= \Im(c_n)$	= 0 pour les fonctions paires
c_n	$= \frac{1}{T} \int_{-T/2}^{T/2} s(t) \cdot e^{-j \cdot 2 \cdot \pi \cdot \frac{n}{T} \cdot t} dt$	$= a_n + j \cdot b_n$	Réel pour les fonctions paires Imaginaire pur pour les fonctions impaires

3. EN PRATIQUE

Nous allons étudier ici deux cas courants de décomposition en série de Fourier : la fonction « carrée » et la fonction « triangle ».

a) Fonction « carrée »

$$s(t) \begin{cases} = 1 & t \in [n \cdot T - T/4, n \cdot T + T/4[\\ = 0 & t \in [n \cdot T + T/4, (n+1) \cdot T - T/4[\end{cases}$$



Avant de développer les calculs, nous pouvons observer les propriétés de la fonction à analyser :

$s(t)$ est paire donc les coefficients b_n sont nuls ;

dans l'intervalle $[-T/2, T/2]$, $s(t)$ est non nulle uniquement dans l'intervalle $[-T/4, T/4]$, il est donc inutile d'effectuer les calculs en dehors de ce dernier.

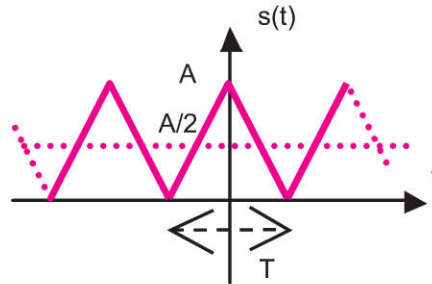
$$a_n = \frac{2}{T} \int_{-T/4}^{T/4} A \cdot \cos\left(n \cdot 2 \cdot \pi \cdot \frac{t}{T}\right) dt$$

$$a_n = A \cdot \frac{2}{T} \cdot \frac{T}{n \cdot 2 \cdot \pi} \cdot \left[\sin\left(n \cdot 2 \cdot \pi \cdot \frac{t}{T}\right) \right]_{-T/4}^{T/4}$$

$$a_n = A \cdot \frac{2 \cdot \sin\left(\frac{n \cdot \pi}{2}\right)}{n \cdot \pi} = A \cdot \text{sinc}\left(\frac{n \cdot \pi}{2}\right)$$

Remarque

La fonction $\text{sinc}(x)$ est la fonction sinus cardinal dont la définition utilisée ici est $\text{sinc}(x) = \frac{\sin(x)}{x}$

b) Fonction triangle

Avant de développer les calculs, nous pouvons observer les propriétés de la fonction à analyser :

$s(t)$ est paire donc les coefficients b_n sont nuls ;

$$a_n = \frac{4 \cdot A}{T} \cdot \left[\int_{-T/4}^{T/4} \left(-\frac{2}{T} \cdot t + 1 \right) \cdot \cos \left(n \cdot 2 \cdot \pi \cdot \frac{t}{T} \right) dt \right]$$

Cette intégrale se calcule en utilisant une intégration par partie dont nous rappelons la formule :

$$\int u(t) \cdot v'(t) dt = [u(t) \cdot v(t)] - \int u'(t) \cdot v(t) dt$$

On pose

$$u(t) = -\frac{2}{T} \cdot t + 1 \Rightarrow u'(t) = -\frac{2}{T}$$

$$v'(t) = \cos \left(n \cdot 2 \cdot \pi \cdot \frac{t}{T} \right) \Rightarrow v(t) = \frac{T}{n \cdot 2 \cdot \pi} \cdot \cos \left(n \cdot 2 \cdot \pi \cdot \frac{t}{T} \right) + K$$

Note

Du fait de la parité de la fonction, nous pouvons calculer les a_n sur l'intervalle $[0, T/2]$ en appliquant un coefficient 2

$$a_n = \frac{4 \cdot A}{T} \cdot \left\{ \left[\left(-\frac{2}{T} \cdot t + 1 \right) \cdot \frac{1}{n \cdot \omega} \sin(n \cdot \omega \cdot t) \right]_0^{T/2} - \int_0^{T/2} -\frac{2}{T} \cdot \frac{1}{n \cdot \omega} \cdot \sin(n \cdot \omega \cdot t) dt \right\} \text{ avec}$$

$$\frac{T}{n \cdot 2 \cdot \pi} = n \cdot \omega$$

On trouve $a_n = \frac{8 \cdot A}{(T \cdot n \cdot \omega)^2} \cdot (-\cos(n \cdot \pi) + 1)$, donc $a_n = 0$ pour n pair et $a_n = \frac{4}{(n \cdot \pi)^2}$ pour n impair.

Il est maintenant possible de reconstituer $s(t)$ à partir de ces coefficients :

$$s(t) = \frac{1}{2} + \frac{4}{\pi^2} \sum_{p=0}^{\infty} \cos[(2 \cdot p + 1) \cdot \omega \cdot t]$$

16 Transformée de Fourier

Mots clés

Spectre, fonction échelon, fonction porte, Dirac, analyse, synthèse.

1. EN QUELQUES MOTS

À l'image d'un prisme décomposant la lumière en un spectre lumineux ou d'un rideau de pluie produisant un arc-en-ciel, la transformée de Fourier permet de représenter tout signal temporel dans le domaine fréquentiel, c'est-à-dire de calculer son **spectre** et de visualiser ainsi sa composition en fréquences.

La décomposition en **séries de Fourier** permet de **décomposer un signal périodique** en une somme pondérée de fonctions sinusoïdales. Cette décomposition correspond à une représentation fréquentielle d'un signal exprimé en fonction du temps. Dans le cas de **signaux non périodiques**, nous avons recours à la **transformée de Fourier** pour accéder à cette **représentation fréquentielle**.

2. CE QU'IL FAUT RETENIR

Quelques fonctions élémentaires très employées en traitement du signal nous seront utiles par la suite : les fonctions échelon et rectangle et l'impulsion (ou distribution) de Dirac ([fiche 14](#)).

a) Transformée de Fourier : analyse

L'analyse au sens de la transformée de Fourier consiste à décomposer un signal intégrable $s(t)$ en une suite continue (en fréquence ou en pulsation) de sinusoïdes pondérées $S(\omega)$.

Rappel

$$e^{j \cdot x} = \cos(x) + j \cdot \sin(x) \Rightarrow e^{-j \cdot x} = \cos(x) - j \cdot \sin(x)$$

$$S(\omega) = \frac{1}{2\pi} \int_{-\infty}^{+\infty} s(t) \cdot e^{-j \cdot \omega \cdot t} dt$$

$$X(n \cdot \omega_0) = c_n \cdot \delta(\omega - n \cdot \omega_0)$$

b) Synthèse

La synthèse est l'opération réciproque de l'analyse : elle consiste, à partir d'une suite continue (en fréquence ou en pulsation) de sinusoïdes pondérées $S(\omega)$, à recomposer un signal temporel $s(t)$.

$$s(t) = \frac{1}{2\pi} \int_{-\infty}^{+\infty} S(\omega) \cdot e^{j \cdot \omega \cdot t} d\omega$$

c) Réciprocité de la transformée de Fourier (TF) et de la transformée de Fourier inverse (TF⁻¹)

$$s(t) \xrightleftharpoons[TF^{-1}]{TF} S(\omega)$$

Propriétés de parité

$s(t)$	$S(\omega)$	$ S(\omega) $	$\varphi_S(\omega)$
Réelle et paire	Réelle et paire	Paire	$0, \pm\pi$
Réelle et impaire	Imaginaire et impaire	Paire	$\pm\pi/2$
Imaginaire et paire	Imaginaire et paire	Paire	$\pm\pi/2$
Imaginaire et impaire	Réelle et impaire	Paire	$0, \pm\pi$
Réelle et quelconque	$\Re(S)$ paire $\Im(S)$ impaire	Paire	Impaire
Imaginaire et quelconque	$\Re(S)$ impaire $\Im(S)$ paire	Paire	Impaire

Tout signal mesuré est réel, il répond donc à la propriété de symétrie Hermitienne :

$$\Re(S) + j \cdot \Im(S) = \Re(-S) - j \cdot \Im(-S)$$

Propriété de similitude

$$x(t) \xrightarrow{TF} S(\omega) \Rightarrow x(a \cdot t) \xrightarrow{TF} \frac{1}{a} \cdot S\left(\frac{\omega}{a}\right)$$

$0 < a < 1$	Étalement de l'échelle des temps	Contraction de l'échelle des fréquences
$1 < a $	Contraction de l'échelle des temps	Étalement de l'échelle des fréquences

Pour que la transformée de Fourier d'un signal existe, il faut que ce signal soit une fonction de carrés sommables, c'est-à-dire que l'intégrale $\int_{-\infty}^{+\infty} s^2(t) dt$ doit être inférieure à l'infini.

L'énergie du signal $s(t)$ doit donc être finie. Cette constatation débouche sur le théorème de Parseval : l'énergie totale W d'un signal répondant à ce critère peut s'exprimer aussi bien à partir de sa représentation dans le domaine complexe que dans le domaine fréquentiel :

$$W_S = \int_{-\infty}^{+\infty} |s(t)|^2 dt = \int_{-\infty}^{+\infty} |S(\omega)|^2 d\omega.$$

3. EN PRATIQUE

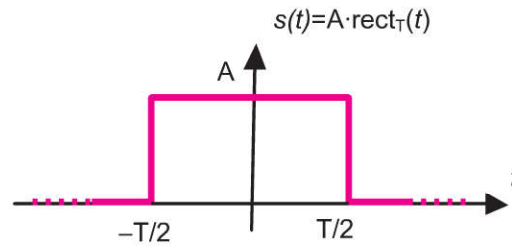
La transformée de Fourier d'un Dirac est 1 :

$$TF[\delta(t)] = \int_{-\infty}^{+\infty} \delta(t) \cdot e^{-j \cdot 2 \cdot \pi \cdot f \cdot t} dt = e^0 = 1$$

a) Autres propriétés relatives au Dirac

$$\begin{aligned} \delta(t - \tau) &\xleftrightarrow{TF} e^{-j \cdot 2 \cdot \pi \cdot \tau \cdot f} & \delta(t + \tau) &\xleftrightarrow{TF} e^{j \cdot 2 \cdot \pi \cdot \tau \cdot f} \\ e^{j \cdot 2 \cdot \pi \cdot f_0 \cdot t} &\xleftrightarrow{TF} \delta(f - f_0) & e^{-j \cdot 2 \cdot \pi \cdot f_0 \cdot t} &\xleftrightarrow{TF} \delta(f + f_0) \end{aligned}$$

b) Transformée de Fourier de la fonction rectangle



Appliquons la formule d'analyse de la transformée de Fourier à la fonction $s(t)$ afin de calculer $S(f)$.

$$S(f) = \int_{-\infty}^{+\infty} s(t) \cdot e^{-j \cdot 2 \cdot \pi \cdot f \cdot t} dt$$

La fonction est bornée, elle est nulle pour des valeurs de t inférieures à $-T/2$, nulle pour t supérieures à $T/2$ et égale à 1 dans cet intervalle.

$$S(f) = A \cdot \int_{-T/2}^{T/2} e^{-j \cdot 2 \cdot \pi \cdot f \cdot t} dt$$

L'étape suivante consiste à effectuer le calcul de cette intégrale

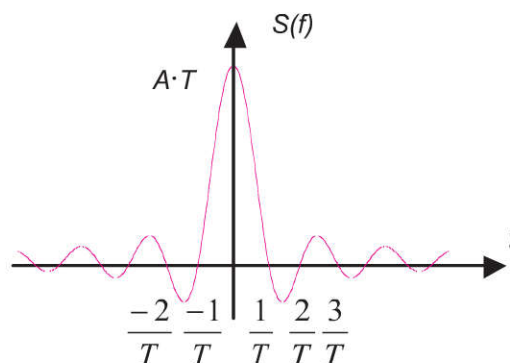
$$\begin{aligned} S(f) &= \frac{A}{-j \cdot 2 \cdot \pi \cdot f} \cdot \left[e^{-j \cdot 2 \cdot \pi \cdot f \cdot t} \right]_{-T/2}^{T/2} \\ &= \frac{A}{\pi \cdot f} \cdot \frac{e^{j \cdot \pi \cdot f \cdot T} - e^{-j \cdot \pi \cdot f \cdot T}}{2 \cdot j} \end{aligned}$$

Sachant que

$$\sin(x) = \frac{e^{j \cdot x} - e^{-j \cdot x}}{2 \cdot j}$$

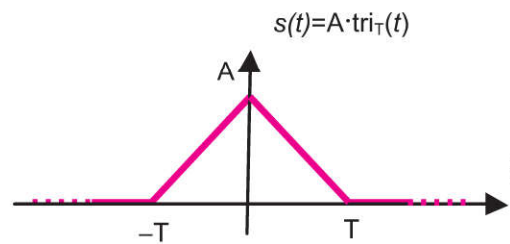
Nous obtenons

$$\begin{aligned} &= \frac{A}{\pi \cdot f} \cdot \sin(\pi \cdot f \cdot T) \\ &= A \cdot T \cdot \operatorname{sinc}(f \cdot T) \end{aligned}$$



c) Transformée de Fourier de la fonction triangle

Le calcul de la transformée de Fourier d'une fonction triangle telle que développée pour la fonction rectangle est fastidieux, c'est pourquoi nous adoptons une démarche différente et nous appliquons certaines propriétés indiquées dans cette fiche.



Une fonction triangle telle qu'elle est décrite sur le diagramme précédent est égale au produit de convolution de deux fonctions rectangle (fiche 18) :

$$s(t) = \text{rect}_T(t) * \text{rect}_T(t)$$

Utilisons la propriété de la transformée de Fourier du produit de convolution pour calculer $S(f)$, transformée de Fourier de la fonction triangle $s(t)$:

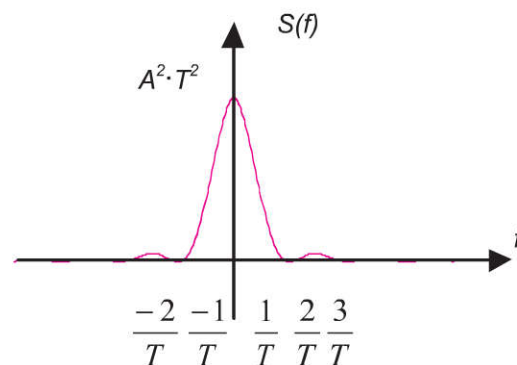
$$S(f) = TF[\text{rect}_T(t) * \text{rect}_T(t)]$$

La transformée de Fourier d'un produit de convolution se décompose en un produit de transformées de Fourier :

$$S(f) = TF[\text{rect}_T(t)] \cdot TF[\text{rect}_T(t)]$$

En reprenant le résultat de l'analyse d'une fonction rectangle, il vient

$$S(f) = [A \cdot T \cdot \text{sinc}(f \cdot T)]^2 = A^2 \cdot T^2 \cdot \text{sinc}^2(f \cdot T)$$



$S(f)$ s'annule pour $f = \frac{k}{T}$ avec $k \in \mathbb{N}$.

17 Transformée de Laplace

Mots clés

Transformée de Fourier, p , impédance opérationnelle, analyse des circuits, équations différentielles.

1. EN QUELQUES MOTS

La transformée de Laplace est très utilisée pour résoudre des équations différentielles et déterminer la fonction de transfert d'un système linéaire car la dérivation et l'intégration dans le domaine de Laplace consistent en de simples divisions et multiplications par la variable p . Elle est similaire à la transformée de Fourier, mais cette dernière ne permet que d'analyser le comportement d'un système en régime établi alors que la transformée de Laplace s'intéresse également au régime transitoire.

2. CE QU'IL FAUT RETENIR

a) Définition

En mathématiques (analyse fonctionnelle), la transformée de Laplace d'une fonction f d'une variable réelle positive t est la fonction F de la variable complexe p , qui est définie par

$$F(p) = \int_0^{+\infty} f(t)e^{-(pt)} dt$$

F est appelée image de f par transformée de Laplace.

p est une variable complexe telle que $p = \sigma + j \cdot \omega$

La transformée de Laplace devient équivalente à la transformée de Fourier lorsque $\sigma = 0$ et donc que $p = j \cdot \omega = j \cdot 2 \cdot \pi \cdot f$

Remarque

Cette formule ne peut-être appliquée qu'à des signaux causaux, c'est-à-dire nuls pour t négatif.

b) Propriétés

	Dilatation	Puissance	Dérivation	Intégration	Exponentielle	Addition	Produit
Fonction	$f(a \cdot t)$	$(-t)^n \cdot f(t)$	$\frac{df(t)}{dt}$	$\int_0^x f(t) dt$	$e^{-\beta \cdot t} \cdot f(t)$	$f(t) + \beta \cdot g(t)$	$f(t) \cdot g(t)$
Transformée	$\frac{1}{a} \cdot F\left(\frac{p}{a}\right)$	$\frac{d^n \cdot F(p)}{dp^n}$	$p \cdot F(p) - f(0^+)$	$\frac{1}{p} \cdot F(p)$	$F(p + \beta)$	$F + \beta \cdot G$	$F(p) * G(p)$

3. EN PRATIQUE

a) Impédances opérationnelles

La **fiche 5** présente les équations relatives au comportement temporel des résistances, condensateurs et inductances. Dans le domaine de Laplace, une impédance se définit comme le rapport entre la tension et le courant :

$$Z(p) = \frac{u(p)}{i(p)}$$

Composant	Résistance	Condensateur	Inductance
Équation différentielle	$u(t) = R \cdot i(t)$	$i(t) = C \cdot \frac{du(t)}{dt}$	$u(t) = L \cdot \frac{di(t)}{dt}$
Impédance opérationnelle	$Z_R(p) = R$ $\Rightarrow v(p) = R \cdot i(p)$	$Z_C(p) = 1/p \cdot C$ $\Rightarrow i(p) = C \cdot p \cdot u(p)$	$Z_L(p) = L \cdot p$ $\Rightarrow u(p) = L \cdot p \cdot i(p)$

Ci-dessous, le tableau regroupe les transformées de Laplace de fonctions usuelles :

Fonction	Transformée	Fonction	Transformée
1	$\frac{1}{p}$	$t \cdot e^{-\beta \cdot t}$	$\frac{1}{(p + \beta)^2}$
t	$\frac{1}{p^2}$	$\frac{t^n \cdot e^{-\beta \cdot t}}{n!}$	$\frac{1}{(p + \beta)^{n+1}}$
$\frac{t^{n-1}}{(n-1)!}$	$\frac{1}{p^n}$	$\sin(\omega \cdot t)$	$\frac{\omega}{p^2 + \omega^2}$
$\delta(t)$	1	$\cos(\omega \cdot t)$	$\frac{p}{p^2 + \omega^2}$
$e^{-\beta \cdot t}$	$\frac{1}{p + \beta}$	$e^{-\beta \cdot t} \cdot \cos(\omega \cdot t)$	$\frac{p + \beta}{(p + \beta)^2 + \omega^2}$

Il suffit de transposer une équation différentielle dans le domaine de Laplace pour obtenir une équation simple à manipuler. Par exemple, la tension aux bornes d'un circuit RL en série est définie par :

$$e(t) = R \cdot i(t) + L \cdot \frac{di(t)}{dt}$$

Dans le domaine de Laplace, cela devient :

$$E(p) = R \cdot I(p) + p \cdot L \cdot I(p)$$

18 Le produit de convolution

Mots clés

Systèmes linéaires, réponse impulsionnelle, causalité, peigne de Dirac.

1. EN QUELQUES MOTS

Le produit de convolution est un calcul intégral impliquant deux fonctions que nous nommerons par la suite $x(t)$ et $h(t)$. Cette formule est très utile pour calculer la sortie d'un système linéaire invariant dans le temps (ou « SLIT » tel qu'un filtre analogique, [fiches 26 et 27](#)) de réponse impulsionnelle $h(t)$ excité par un signal d'entrée $x(t)$ quelconque.

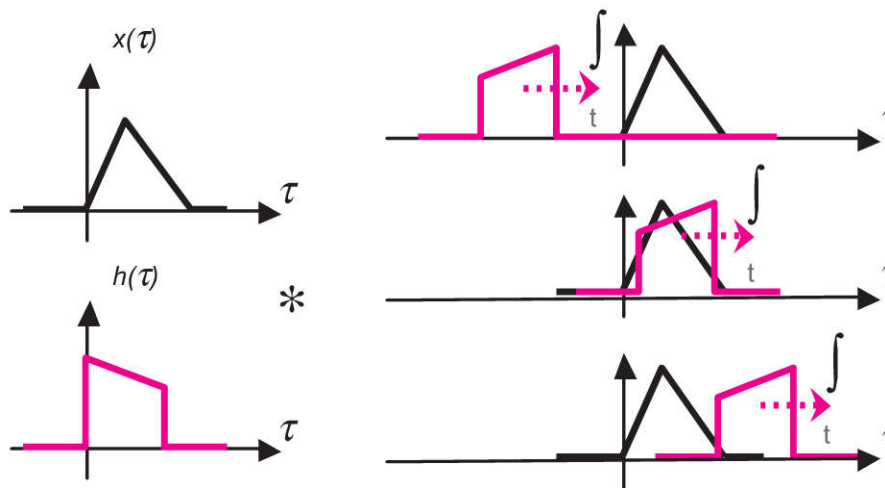
2. CE QU'IL FAUT RETENIR

a) Définitions

Soient $x(\tau)$ et $h(\tau)$ deux signaux causaux et $y(t)$ le résultat du produit de convolution de x et de h .

$$y(t) = x(t) * h(t) = \int_0^t x(\tau) \cdot h(t - \tau) d\tau$$

Cela revient à considérer que la fonction $x(t)$ est fixe et que la fonction $h(-t)$ glisse sur l'axe des temps.

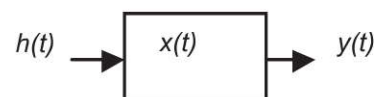
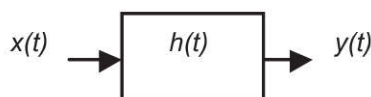


Remarque

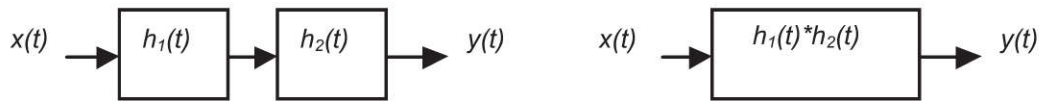
Un signal est causal s'il vaut zéro pour $t < 0$, c'est-à-dire qu'il n'existe pas de signal à proprement parler avant l'origine des temps.

b) Propriétés

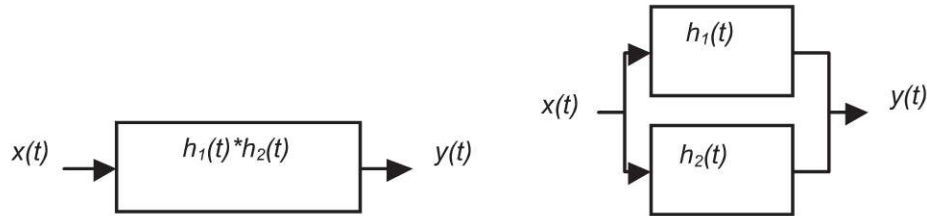
Commutativité : $x(t) * h(t) = h(t) * x(t)$



Associativité : $x(t) * [h_1(t) * h_2(t)] = [x(t) * h_1(t)] * h_2(t)$



Distributivité : $x(t) * [h_1(t) + h_2(t)] = x(t) * h_1(t) + x(t) * h_2(t)$



3. EN PRATIQUE

La convolution d'un signal par une impulsion de Dirac est égale au signal lui-même : le Dirac est l'élément neutre de la convolution.

$$x(t) * \delta(t) = x(t)$$

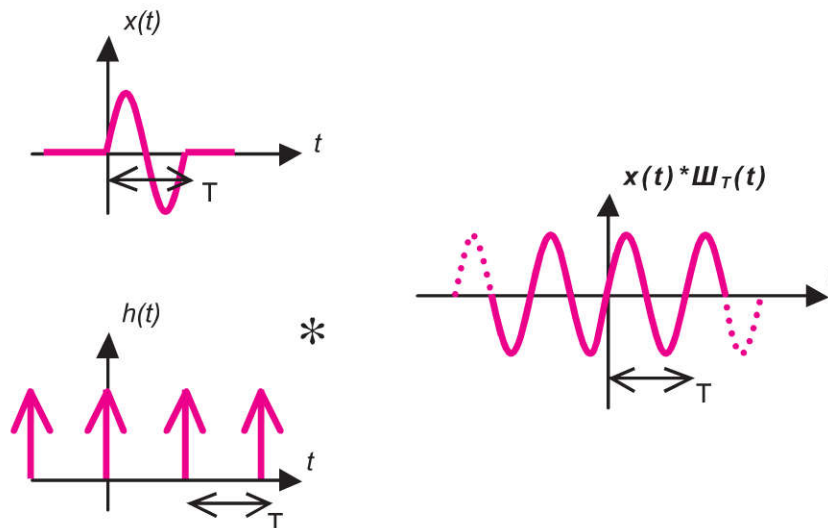
La convolution d'un signal par une impulsion de Dirac décalée dans le temps revient à recopier le signal d'origine avec le même décalage temporel.

$$x(t) * \delta(t - t_0) = x(t - t_0)$$

En combinant cette propriété avec la distributivité du produit de convolution, nous pouvons déduire

$$x(t) * \text{III}_T(t) = \sum_{k=-\infty}^{+\infty} x(t - k \cdot T)$$

Toute fonction périodique peut être décomposée en un produit de convolution d'une période du signal par un peigne de Dirac.



Attention

Il ne faut pas confondre produit de convolution par un Dirac et multiplication par un Dirac, car nous aurions alors :

$$x(t) \cdot \delta(t) = x(0) \cdot \delta(t) \text{ et } x(t) \cdot \delta(t - t_0) = x(t_0) \cdot \delta(t - t_0)$$

19 Échelle logarithmique et décibels

Mots clés

Diagramme de Bode, dB, dBm, dBW, puissance, quadripôle, filtrage, additivité des gains.

1. EN QUELQUES MOTS

Il est souvent difficile voire impossible de représenter graphiquement la variation d'une quantité s'étalant sur plusieurs ordres de grandeur. La manipulation, à l'écrit ou à l'oral, de puissances de dix élevées et la comparaison de grands nombres peut être compliquée. Les échelles logarithmiques et les dB permettent de palier à ces problèmes en convertissant des grandeurs associées à des puissances de 10 élevées à des nombres plus réduits et de les comparer entre eux.

2. CE QU'IL FAUT RETENIR

a) Comparaison de puissances

La comparaison des puissances entre l'entrée et la sortie d'un quadripôle peut être représentée par un rapport dont le résultat est sans unité. Le gain en dB répond à la formulation suivante :

$$G_{dB} = 10 \cdot \log(A) = 10 \cdot \log\left(\frac{P_s}{P_e}\right) = 10 \cdot \log(P_s) - 10 \cdot \log(P_e)$$

Si la puissance en sortie est supérieure à la puissance d'entrée, A est supérieur à 1 et le gain est supérieur à 0 dB. Réciproquement, une puissance de sortie supérieure à la puissance d'entrée conduit à un gain linéaire A inférieur à 1 et à un gain inférieur à 0 dB.

b) Comparaison de tensions et de courants

Afin d'assurer l'homogénéité de l'échelle dB lorsque l'on compare des puissances ou des tensions, il convient de garder à l'esprit qu'une puissance électrique est proportionnelle au carré d'une tension ou d'un courant.

Voici le calcul permettant de déduire le gain en dB d'un rapport de courant en partant de la puissance dissipée dans une résistance R parcourue par un courant i .

$$P = R \cdot i_{eff}^2 = \frac{1}{2} \cdot R \cdot \hat{i}^2 \Rightarrow G_{dB} = 10 \cdot \log\left(\frac{R \cdot i_s^2}{R \cdot i_e^2}\right) = 10 \cdot \log\left(\frac{i_s}{i_e}\right)^2 = 20 \cdot \log\left(\frac{i_s}{i_e}\right) = 20 \cdot \log\left(\frac{\hat{i}_s}{\hat{i}_e}\right)$$

Le calcul suivant est similaire et permet de déduire le gain en dB d'un rapport de tensions en partant de la puissance dissipée dans une résistance R soumise à une tension v à ses bornes.

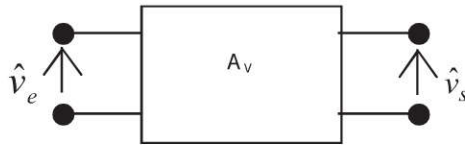
$$P = \frac{v_{eff}^2}{R} = \frac{1}{2} \cdot \frac{\hat{v}^2}{R} \Rightarrow G_{dB} = 10 \cdot \log\left(\frac{\frac{v_s^2}{R}}{\frac{v_e^2}{R}}\right) = 10 \cdot \log\left(\frac{v_s}{v_e}\right)^2 = 20 \cdot \log\left(\frac{v_s}{v_e}\right) = 20 \cdot \log\left(\frac{\hat{v}_s}{\hat{v}_e}\right)$$

Gain linéaire en tension

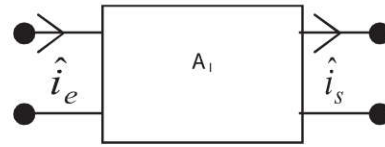
$$A_v = \frac{\hat{v}_s}{\hat{v}_e} = \frac{v_s}{v_e}$$

Gain linéaire en courant

$$A_i = \frac{\hat{i}_s}{\hat{i}_e} = \frac{i_s}{i_e}$$



$$G_{dB}^v = 20 \cdot \log\left(\frac{\hat{v}_s}{\hat{v}_e}\right) = 20 \cdot \log(A_v)$$

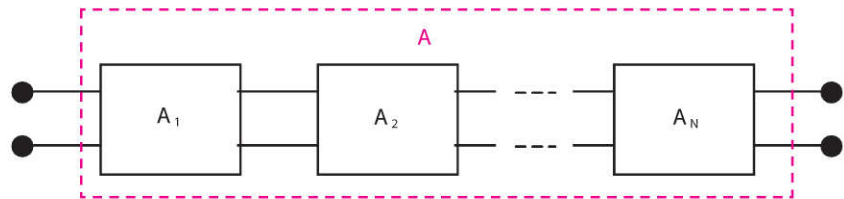


$$G_{dB}^i = 20 \cdot \log\left(\frac{\hat{i}_s}{\hat{i}_e}\right) = 20 \cdot \log(A_i)$$

c) Propriétés

Additivité des gains

En considérant des gains linéaires quels qu'ils soient (tension, courant, puissance), le gain total de N quadripôles mis en cascade est la multiplication des gains de chaque élé-



ment, soit : $A = A_1 \cdot A_2 \cdot \dots \cdot A_N = \prod_{i=1}^N A_i$

En dB, les propriétés des fonctions logarithmiques conduisent à additionner les gains :

$$G_{dB} = G_1 + G_2 + \dots + G_N = \sum_{i=1}^N G_i$$

Rappel : $\log(a \cdot b) = \log(a) + \log(b)$

Remarque

La mise en cascade d'éléments suppose la prise en compte de l'adaptation d'impédance vue dans les [fiches 20](#) et [21](#).

Valeur en dB	Valeurs linéaires		
	Puissance $\frac{P_s}{P_e} = 10^{\frac{G_{dB}}{10}}$	Tension $A_v = \frac{v_s}{v_e} = 10^{\frac{G_{dB}}{20}}$	Courant $A_i = \frac{i_s}{i_e} = 10^{\frac{G_{dB}}{20}}$
-20 dB	0,01	0,1	
-10 dB	0,1	$1/\sqrt{10}$	
-6 dB	0,25	0,5	
-3 dB	0,5	$1/\sqrt{2}$	
0 dB	1	1	
3 dB	2	$\sqrt{2}$	
6 dB	4	2	
10 dB	10	$\sqrt{10}$	
20 dB	100	10	

d) Valeurs de puissance

Le dB ne s'applique que pour représenter des rapports sans dimensions. Cependant, dans certains cas, on utilise ce formalisme pour indiquer des valeurs de puissance : elles sont alors divisées par 1 W et la valeur est notée en dBW (« décibels-watt » ou « dB-watt ») ou divisées par 1 mW, ce qui correspond à l'unité dBm (« décibels-milliwatts »).

$$P_{dBW} = 10 \cdot \log\left(\frac{P_W}{1W}\right) \text{ et } P_{dBm} = 10 \cdot \log\left(\frac{P_{mW}}{1mW}\right)$$

3. EN PRATIQUE

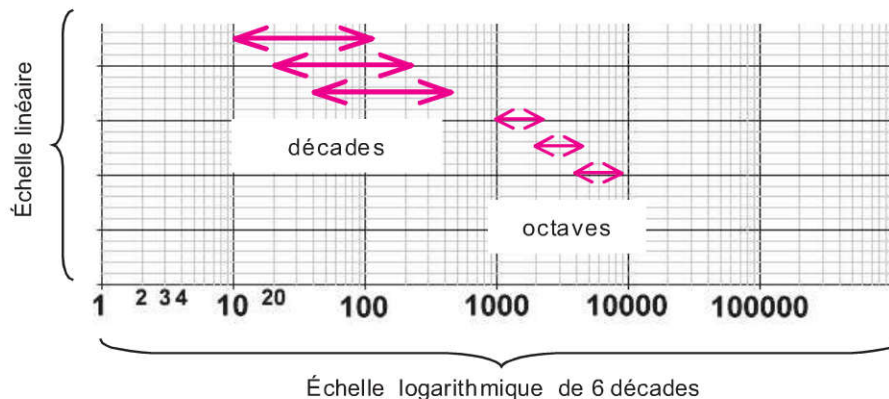
a) Diagramme de Bode

L'axe des fréquences utilise une échelle logarithmique afin de représenter aisément plusieurs décades.

Remarques

- Une **décade** est l'intervalle de fréquences située entre f_0 et $f_1 = 10 \cdot f_0$, elle correspond à l'intervalle entre deux graduations principales successives de l'axe des abscisses.
- Une **octave** est l'intervalle de fréquences située entre f_0 et $f_2 = 2 \cdot f_0$ elle correspond à l'intervalle entre une graduation principale et la graduation secondaire suivante de l'axe des abscisses.

Le papier semi-logarithmique propose une échelle logarithmique que nous associons dans notre domaine aux fréquences et une échelle linéaire qui permet de représenter des valeurs de gain en dB ou de phase (typiquement en degrés ou radians).



Le **diagramme de Bode** est très largement employé pour visualiser le comportement de SLIT (fiche 22) et en particulier leur transmittance (fonction de transfert) dans le domaine fréquentiel. Il se décompose en deux parties : un diagramme de gain et un diagramme de déphasage.

Le gain en dB est obtenu en calculant le module de la transmittance :

$$G_{dB}(\omega) = 20 \cdot \log(|T(\omega)|)$$

Le déphasage est obtenu en calculant l'argument de la transmittance $\varphi(\omega) = \text{Arg}(T(\omega))$

b) Représentation fréquentielle des filtres

Exemple détaillé : circuit RC série avec, tension prise aux bornes du condensateur (fiche 26).

$$\text{La transmittance complexe s'écrit } T(\omega) = \frac{1/j \cdot C \cdot \omega}{R + 1/j \cdot C \cdot \omega} = \frac{1}{1 + j \cdot R \cdot C \cdot \omega}$$

Le module de cette transmittance répond à la formulation suivante : $|T(\omega)| = \frac{1}{\sqrt{1 + R^2 \cdot C^2 \cdot \omega^2}}$

Les gains étant exprimés en dB dans les diagrammes de Bode, nous nous intéressons au module en dB de la transmittance représentant l'atténuation imposée par le filtre en fonction

de la pulsation : $|T(\omega)|_{dB} = 20 \cdot \log\left(\frac{1}{\sqrt{1 + R^2 \cdot C^2 \cdot \omega^2}}\right) = -10 \cdot \log(1 + R^2 \cdot C^2 \cdot \omega^2)$

et à son argument qui représente le déphasage entre la tension d'excitation et la tension de sortie $\text{Arg}(T(\omega)) = \text{Arctan}(R \cdot C \cdot \omega)$

Le fréquence de coupure d'un filtre est définie telle que $|T(\omega_0)| = 1/\sqrt{2}$. On peut donc en

déduire la pulsation ω_0 correspondante : $|T(\omega_0)| = \frac{1}{\sqrt{1 + R^2 \cdot C^2 \cdot \omega_0^2}} = \frac{1}{\sqrt{2}} \Rightarrow \omega_0 = \frac{1}{R \cdot C}$

Le tableau ci-après indique les points caractéristiques des fonctions $|T(\omega)|$ et $\text{Arg}(T(\omega))$, et le comportement aux limites :

	Atténuation	Déphasage
$\omega = 0 \text{ rad} \cdot \text{s}^{-1}$	$ T(0) = 1$	$\text{Arg}(T(0)) = \text{Arctan}(0) = 0$
$\omega = \omega_0$	$ T(\omega_0) = \frac{1}{\sqrt{1 + R^2 \cdot C^2 \cdot \omega_0^2}} = \frac{1}{\sqrt{2}} \Rightarrow \omega_0 = \frac{1}{R \cdot C}$	$\text{Arg}(T(\omega_0)) = \text{Arctan}(1) = -\pi/4$
$\omega \rightarrow \infty$	$\lim_{\omega \rightarrow +\infty} T(\omega) = 0$	$\lim_{\omega \rightarrow +\infty} \text{Arg}(T(\omega)) = -\pi/2$

Les valeurs de gain indiquées dans le tableau nous permettent de déduire :

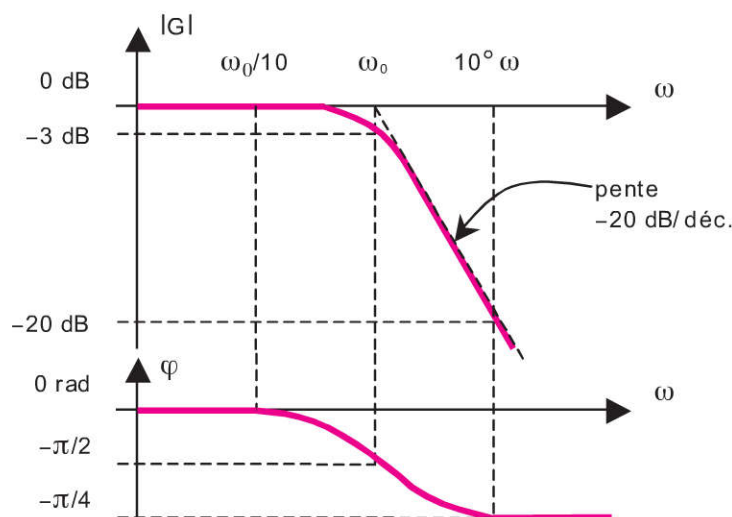
- qu'un filtre passif n'amplifie pas puisque le gain maximal est de 1 ;
- le filtre RC série étudié ici est un filtre passe-bas car T diminue lorsque la pulsation augmente. Pour un filtre passe-bas, ce sont les fréquences les plus basses qui sont le moins atténuées.

Quel est le comportement asymptotique du gain ? Soit une pulsation $\omega_1 \gg \omega_0$ et une pulsa-

tion $\omega_2 = 10 \cdot \omega_1$. Nous utilisons l'approximation suivante : $1 + R^2 \cdot C^2 \cdot \omega_1^2 \approx R^2 \cdot C^2 \cdot \omega_1^2$

La différence de gain entre les deux pulsations distantes d'une décade est :

$$|T(\omega_2)|_{dB} - |T(\omega_1)|_{dB} \approx -20 \text{ dB}$$



20 Adaptation d'impédance

Mots clés

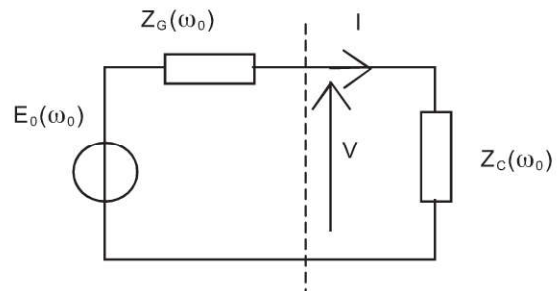
Adaptation en puissance, en courant, en tension, 50 Ohms, collecteur commun, AOP.

1. EN QUELQUES MOTS

Une source réelle indépendante de tension présente une certaine impédance interne. Celle-ci influence le passage du signal en entrée du circuit alimenté par le générateur : en effet, l'impédance interne du générateur mise en série avec l'impédance de charge constitue un filtre (fiches 26 et 27). En fonction du rapport entre ces deux impédances, le passage du courant, de la tension ou de la puissance est favorisé.

2. CE QU'IL FAUT RETENIR

La puissance disponible d'une source indépendante est la puissance maximum $\hat{P}_C$ qu'une source peut délivrer à une charge placée à ses bornes. Soit un générateur de tension E_0 d'impédance interne fixe Z_G et une charge d'impédance Z_C que nous cherchons à optimiser de sorte que la puissance délivrée par le générateur soit maximum.



L'analyse de ce circuit prend la forme d'un mon-

tage diviseur de tension (fiche 23). Soit I le courant qui circule dans la branche $I = \frac{E_0}{Z_G + Z_C}$

avec $Z_C = R_C + j \cdot X_C$ et $Z_G = R_G + j \cdot X_G$

Nous pouvons en déduire la tension V aux bornes de l'impédance de charge :

$$V = Z_C \cdot I = \frac{Z_C \cdot E_0}{Z_G + Z_C}$$

La puissance consommée par la charge est calculée en fonction du courant et de la tension (fiche 12) :

$$P_C = \frac{1}{2} \cdot \Re(V \cdot I^*) = \frac{1}{2} \cdot \frac{R_C \cdot |E_0|^2}{|Z_G + Z_C|^2} = \frac{1}{2} \cdot \frac{R_C \cdot |E_0|^2}{(R_G + R_C)^2 - (X_G + X_C)^2}$$

Selon la réactance X , la puissance est maximum pour $X_G = -X_C$.

Dérivons cette expression de la puissance avec $X_G + X_C = 0$ selon la résistance de charge afin de trouver la valeur de R_C pour laquelle P_C admet un extremum :

$$\left. \frac{dP_C}{dR_C} \right|_{X_G = -X_C} = \frac{|E_0|^2}{2} \cdot \frac{2 \cdot (R_C + R_G) \cdot R_C - (R_C + R_G)^2}{(R_C + R_G)^4} = \frac{|E_0|^2}{2} \cdot \frac{R_C - R_G}{(R_C + R_G)^4}$$

$\frac{dP_C}{dR_C}$ s'annule pour $R_C^2 - R_G^2 = 0$ ce qui implique $R_C = \pm R_G$ sachant que $-R_G$ (résistance négative) est une solution mathématique mais non physique.

Ce calcul démontre donc que la puissance délivrée est maximale lorsque la charge présente une impédance dont la valeur complexe est le conjugué de l'impédance du générateur soit $Z_C = Z_G^*$.

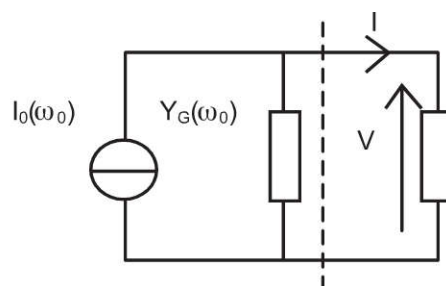
$$\hat{P}_C = \frac{1}{8} \cdot \frac{|E_0|^2}{\Re(Z_G)}$$

Un calcul similaire avec un générateur de courant donnerait lieu à une conclusion similaire :

en considérant les admittances $Y_C = G_C + j \cdot B_C$ et

$$Y_G = G_G + j \cdot B_G.$$

La puissance délivrée est maximale lorsque la charge présente une admittance dont la valeur complexe est le conjugué de l'admittance du générateur soit $Y_C = Y_G^*$.



Cette puissance s'écrit : $\hat{P}_C = \frac{1}{8} \cdot \frac{|I_0|^2}{\Re(Y_G)}$

En reprenant les deux schémas précédents, nous pouvons facilement optimiser le rapport des impédances (ou admittances) pour que le courant ou la tension au niveau de la charge soit maximum.

Générateur de tension	Adaptation en tension	$V_E = \frac{Z_E}{Z_G + Z_E} \cdot V_G$	$Z_E \rightarrow \infty \Rightarrow V_E \rightarrow V_G$	Circuit ouvert
	Adaptation en courant	$I = \frac{V_G}{Z_G + Z_E}$	$Z_E = 0 \Rightarrow I = \frac{V_G}{Z_G}$	Court-circuit
Générateur de courant	Adaptation en tension	$U = \frac{I_G}{Y_G + Y_E}$	$Y_E = 0 \Rightarrow U = \frac{I_G}{Y_G}$	Circuit ouvert
	Adaptation en courant	$I_E = \frac{Y_E}{Y_G + Y_E} \cdot I_G$	$Y_E \rightarrow \infty \Rightarrow I_E \rightarrow I_G$	Court-circuit

3. EN PRATIQUE

Les montages adaptateurs d'impédance permettent d'optimiser le courant, la tension ou la puissance délivrée par un générateur à une charge tel que le montage collecteur commun (fiche 38).

Les meilleures impédances de câbles coaxiaux ont été déterminées par les laboratoires Bell en 1929. La valeur de 50 Ohms retenue en général et correspondant à l'impédance interne standard pour un générateur correspond à un bon compromis pour pouvoir véhiculer des puissances élevées, des hautes tensions avec une faible atténuation.

21 Quadripôles

Mots clés

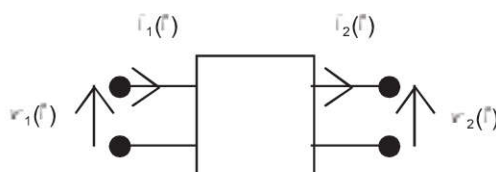
Matrice chaîne, matrice impédance, matrice admittance, mise en cascade.

1. EN QUELQUES MOTS

La mise en cascade de différents éléments d'un circuit tels que des amplificateurs, des adaptateurs d'impédance ou des filtres peut conduire à une analyse fastidieuse. En considérant ces éléments comme des quadripôles et en appliquant un formalisme matriciel, il est possible de simplifier l'interprétation, la modélisation et l'analyse d'un système électronique.

2. CE QU'IL FAUT RETENIR

Un quadripôle présente deux accès.



Pour un quadripôle, il existe quatre variables : $v_1(t)$, $v_2(t)$, $i_1(t)$ et $i_2(t)$. Deux de ces variables sont considérées comme des excitations (entrées) et deux comme des réponses (sorties). Quatre variables et deux accès laissent six façons primaires de décrire un dipôle sous forme matricielle.

► Matrice hybride

$$\begin{bmatrix} v_1(t) \\ i_2(t) \end{bmatrix} = [h(t)] * \begin{bmatrix} i_1(t) \\ v_2(t) \end{bmatrix} \text{ avec } [h(t)] = \begin{bmatrix} h_{11}(t) & h_{12}(t) \\ h_{21}(t) & h_{22}(t) \end{bmatrix}$$

► Matrice impédance

$$\begin{bmatrix} v_1(t) \\ v_2(t) \end{bmatrix} = \begin{bmatrix} z_{11}(t) & z_{12}(t) \\ z_{21}(t) & z_{22}(t) \end{bmatrix} * \begin{bmatrix} i_1(t) \\ i_2(t) \end{bmatrix} = [z(t)] * \begin{bmatrix} i_1(t) \\ i_2(t) \end{bmatrix}$$

► Matrice d'admittance

$$\begin{bmatrix} i_1(t) \\ i_2(t) \end{bmatrix} = [y(t)] * \begin{bmatrix} v_1(t) \\ v_2(t) \end{bmatrix}$$

► Matrice chaîne

$$\begin{bmatrix} v_1(t) \\ i_1(t) \end{bmatrix} = [c(t)] * \begin{bmatrix} v_2(t) \\ -i_2(t) \end{bmatrix}$$

Par transformée de Laplace ([fiche 17](#)), la **fonction d'excitation en courant** (matrice impédance) devient :

$$\begin{bmatrix} V_1(p) \\ V_2(p) \end{bmatrix} = \begin{bmatrix} Z_{11}(p) & Z_{12}(p) \\ Z_{21}(p) & Z_{22}(p) \end{bmatrix} * \begin{bmatrix} I_1(p) \\ I_2(p) \end{bmatrix}$$

Z_{11} : **impédance d'entrée** du quadripôle

Z_{12} : **impédance de transfert inverse** du quadripôle

Z_{21} : **impédance de transfert** du quadripôle

Z_{22} : **impédance de sortie** du quadripôle.

Par **transformée de Laplace**, la relation tension-courant utilisant la **matrice admittance** devient :

$$\begin{bmatrix} I_1(p) \\ I_2(p) \end{bmatrix} = \begin{bmatrix} Y_{11}(p) & Y_{12}(p) \\ Y_{21}(p) & Y_{22}(p) \end{bmatrix} * \begin{bmatrix} V_1(p) \\ V_2(p) \end{bmatrix}$$

Y_{11} : **admittance d'entrée** du quadripôle

Y_{12} : **admittance de transfert inverse** du quadripôle

Y_{21} : **admittance de transfert** du quadripôle

Y_{22} : **admittance de sortie** du quadripôle.

Remarque

Les matrices $[Y]$ et $[Z]$ sont inverses l'une de l'autre :

$$[U] = [Z] \cdot [I] \text{ et } [I] = [Y] \cdot [U] \text{ donc } [U] = [Z] \cdot [Y] \cdot [U]. \text{ On en déduit que } [Z] = [Y]^{-1}.$$

Les **paramètres ABCD** sont aussi connus sous le nom de paramètre chaîne, cascade ou de ligne de transmission ; ils s'écrivent de la manière suivante après transformée de Laplace :

$$\begin{bmatrix} V_1(p) \\ I_1(p) \end{bmatrix} = \begin{bmatrix} A(p) & B(p) \\ C(p) & D(p) \end{bmatrix} * \begin{bmatrix} V_2(p) \\ -I_2(p) \end{bmatrix}$$

Pour des problèmes de mesure en radiofréquences et microondes, les quatre variables primaires $v_1(t)$, $v_2(t)$, $i_1(t)$ et $i_2(t)$ peuvent être combinées afin de former de nouvelles excitations et réponses : les ondes de puissance. À la différence des matrices $[Y]$ et $[Z]$, les **paramètres S** ne se basent pas sur des valeurs de signal (tension, courant) ni sur des calculs en circuit ouvert ou court circuit mais sur des valeurs homogènes à des puissances et sur l'adaptation ou non des impédances. Les paramètres S sont définis en termes d'ondes incidentes et réfléchies sur les ports.

Soient $\alpha = \frac{1}{2 \cdot \sqrt{R_0}}$ et $\beta = \frac{\sqrt{R_0}}{2}$. Ces constantes permettent de définir les excitations et les réponses suivantes :

$$a_1(t) = \alpha \cdot v_1(t) + \beta \cdot i_1(t) \text{ et } a_2(t) = \alpha \cdot v_2(t) + \beta \cdot i_2(t)$$

$$b_1(t) = \alpha \cdot v_1(t) - \beta \cdot i_1(t) \text{ et } b_2(t) = \alpha \cdot v_2(t) - \beta \cdot i_2(t)$$

On écrit sous forme matricielle

$$\begin{bmatrix} b_1(t) \\ b_2(t) \end{bmatrix} = \begin{bmatrix} S_{11}(t) & S_{12}(t) \\ S_{21}(t) & S_{22}(t) \end{bmatrix} * \begin{bmatrix} a_1(t) \\ a_2(t) \end{bmatrix}$$

Ce qui peut s'écrire après transformée de Laplace sous la forme

$$\begin{bmatrix} b_1(p) \\ b_2(p) \end{bmatrix} = \begin{bmatrix} S_{11}(p) & S_{12}(p) \\ S_{21}(p) & S_{22}(p) \end{bmatrix} \cdot \begin{bmatrix} a_1(p) \\ a_2(p) \end{bmatrix}$$

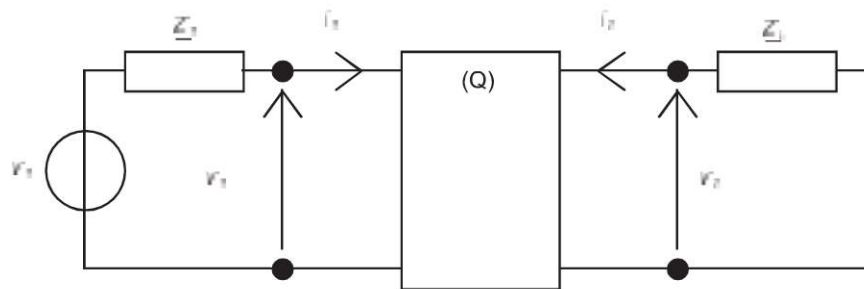
Le rapport $\rho = \frac{b(p)}{a(p)}$ définit le coefficient de réflexion.

Tous les circuits n'ont pas de matrice $[Z]$ ou $[Y]$, mais ils ont tous une matrice chaîne et une matrice $[S]$.

3. EN PRATIQUE

a) Calcul des matrices impédance et admittance

Soit le quadripôle Q suivant excité par un générateur de tension v_1 et d'impédance interne Z_1 et refermé sur une impédance Z_L



$$[U] = [Z] \cdot [I] \Leftrightarrow \begin{cases} u_1 = z_{11} \cdot i_1 + z_{12} \cdot i_2 \\ u_2 = z_{21} \cdot i_1 + z_{22} \cdot i_2 \end{cases} \quad [I] = [Y] \cdot [U] \Leftrightarrow \begin{cases} i_1 = y_{11} \cdot v_1 + y_{12} \cdot v_2 \\ i_2 = y_{21} \cdot v_1 + y_{22} \cdot v_2 \end{cases}$$

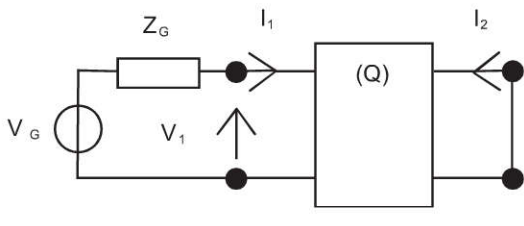
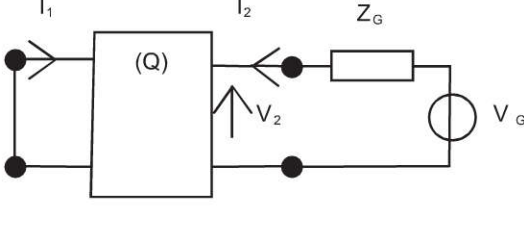
Les paramètres z_{1i} sont calculés en considérant que le courant i_2 est nul, c'est-à-dire que le quadripôle est en circuit ouvert en sortie.

Les paramètres z_{2i} sont calculés en considérant que le courant i_1 est nul, c'est-à-dire que le quadripôle est en circuit ouvert en entrée.

$z_{11} = \frac{u_1}{i_1} \Big _{i_2=0}$	
$z_{12} = \frac{u_1}{i_2} \Big _{i_1=0}$	
$z_{21} = \frac{u_2}{i_1} \Big _{i_2=0}$	
$z_{22} = \frac{u_2}{i_2} \Big _{i_1=0}$	

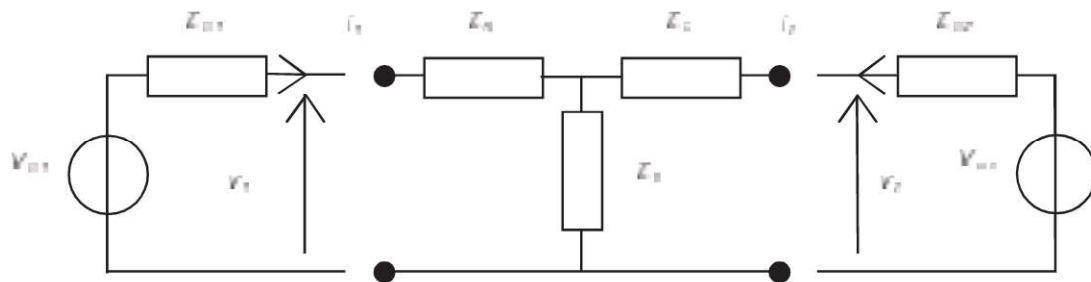
Les paramètres y_{1i} sont calculés en considérant que la tension v_2 est nulle, c'est-à-dire que le quadripôle est en court-circuit en sortie.

Les paramètres y_{2i} sont calculés en considérant que la tension v_1 est nulle, c'est-à-dire que le quadripôle est en court-circuit en entrée.

$y_{11} = \frac{i_1}{u_1} \Big _{v_2=0}$	
$y_{12} = \frac{i_1}{v_2} \Big _{v_1=0}$	
$y_{21} = \frac{i_1}{v_1} \Big _{v_2=0}$	
$y_{22} = \frac{i_2}{v_2} \Big _{v_1=0}$	

b) Exemples

L'exemple suivant montre les résultats de calcul pour la matrice impédance de trois impédances placées en T conformément au schéma suivant :



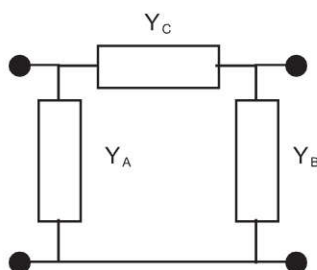
$$z_{11} = \frac{u_1}{i_1} \Big|_{i_2=0} = Z_A + Z_B \quad z_{12} = \frac{u_1}{i_2} \Big|_{i_1=0} = Z_B$$

$$z_{21} = \frac{u_2}{i_1} \Big|_{i_2=0} = Z_B \quad z_{22} = \frac{u_2}{i_2} \Big|_{i_1=0} = Z_B + Z_C$$

$$[Z] = \begin{bmatrix} Z_A + Z_B & Z_B \\ Z_B & Z_B + Z_C \end{bmatrix}$$

La matrice admittance du quadripôle ci-joint s'écrit de la manière suivante :

$$[Y] = \begin{bmatrix} Y_A + Y_B & Y_B \\ Y_B & Y_B + Y_C \end{bmatrix}$$



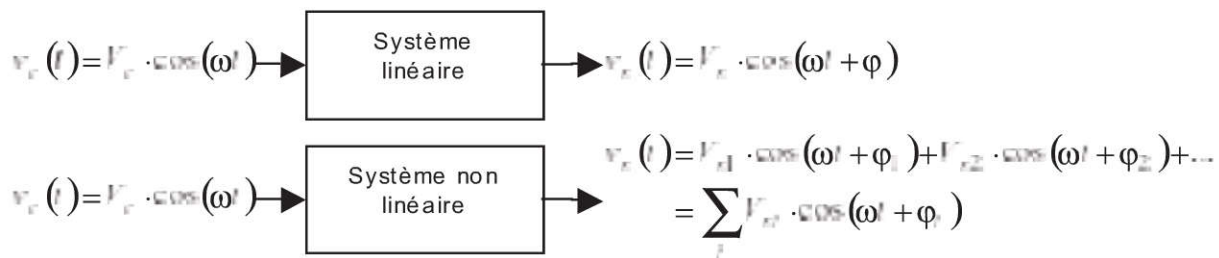
22 Systèmes linéaires invariants dans le temps

Mots clés

Réponse impulsionnelle, convolution, symétrie hermitienne, diagramme de Nyquist.

1. EN QUELQUES MOTS

Un système linéaire atténue, amplifie et éventuellement déphase le signal sinusoïdal qui lui est appliqué en entrée mais il ne le déforme pas.



2. CE QU'IL FAUT RETENIR

a) Produit de convolution

En mathématiques, le produit de convolution de deux fonctions réelles ou complexes f et g se note généralement « $*$ » et s'écrit (fiche 18) :

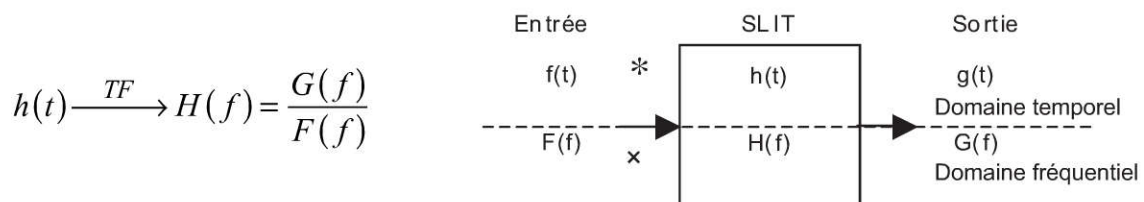
$$(e * h)(t) = \int_{-\infty}^{+\infty} s(t - \tau) \cdot h(\tau) \cdot d\tau = \int_{-\infty}^{+\infty} s(\tau) \cdot h(t - \tau) \cdot d\tau$$

Un système linéaire invariant dans le temps (SLIT) effectue une convolution : la sortie $s(t)$ est égale au produit de convolution de l'entrée $f(t)$ par la **réponse impulsionnelle $h(t)$ du SLIT**.

$$\left. \begin{array}{l} f(t) \xrightarrow{TF} F(f) \\ h(t) \xrightarrow{TF} H(f) \end{array} \right\} s(t) = f(t) * h(t) \xrightarrow{TF} S(f) = F(f) \cdot H(f)$$

$\xleftarrow{TF^{-1}}$

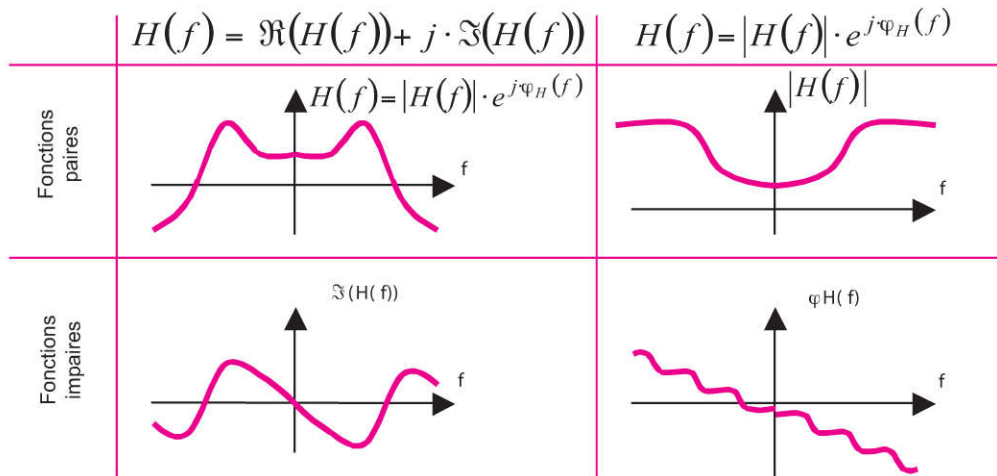
On peut en déduire que la réponse en fréquence d'un SLIT est le rapport entre le spectre de sortie et le spectre d'entrée.



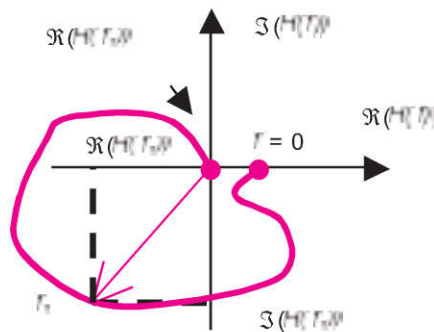
Dans le cas général, $H(f)$ est complexe :

$H(f) = |H(f)| \cdot e^{j \cdot \varphi_H(f)}$ avec $|H(f)|$ le spectre en amplitude du système et $\varphi_H(f)$ son spectre de phase. Dans le cas des **systèmes réels**, les **propriétés de symétrie hermitienne** s'appliquent à $H(f)$:

$$\begin{aligned} \Re(H(f)) \text{ paire} &\Rightarrow \Re(H(f)) = \Re(H(-f)) & |H(f)| \text{ paire} &\Rightarrow |H(f)| = |H(-f)| \\ \Im(H(f)) \text{ impaire} &\Rightarrow \Im(H(f)) = -\Im(H(-f)) & \varphi_H(f) \text{ impaire} &\Rightarrow \varphi_H(f) = -\varphi_H(-f) \end{aligned}$$

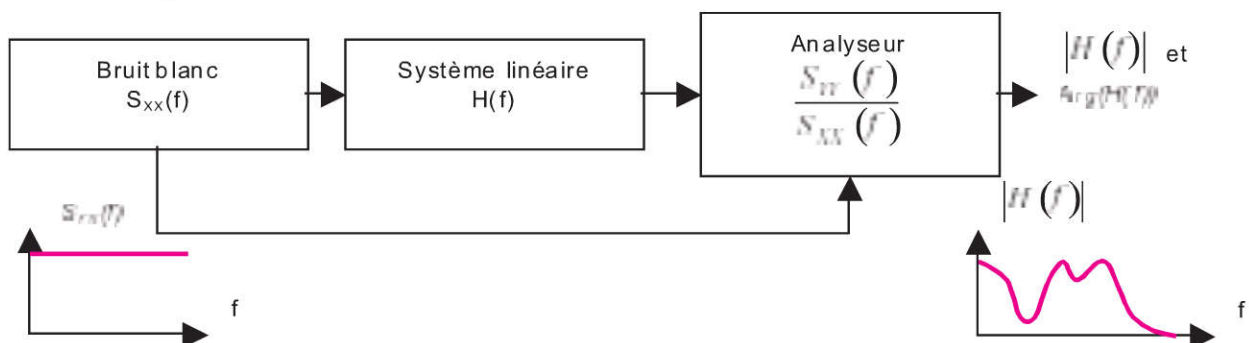


Une représentation très utilisée pour visualiser H dans le domaine fréquence est le diagramme de Nyquist : l'axe des abscisses correspond à la partie réelle de la fonction de transfert du SLIT et l'axe des ordonnées à sa partie imaginaire.



3. EN PRATIQUE

La caractérisation d'un système linéaire peut s'effectuer à l'aide d'un bruit blanc ou d'une impulsion. La [fiche 68](#) indique qu'un bruit blanc est un signal aléatoire présentant une densité spectrale de puissance constante quelle que soit la fréquence, en d'autres termes, un bruit blanc contient toutes les fréquences avec la même puissance. Utiliser un bruit blanc en entrée d'un système linéaire permet d'exciter celui-ci sur une bande de fréquence infinie et de mesurer la réponse en fréquence $|H(f)|$, c'est-à-dire l'atténuation ou l'amplification en sortie en fonction de la fréquence.



23 Montage parallèle, montage série

Mots clés

Diviseur de courant, diviseur de tension, composants passifs.

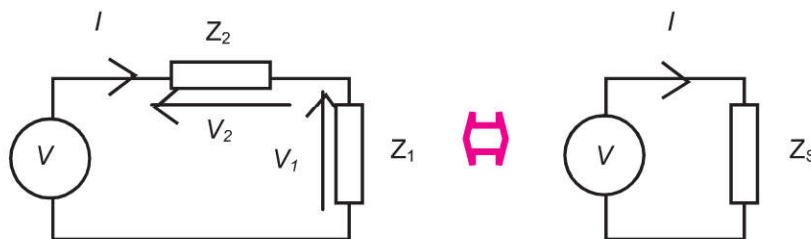
1. EN QUELQUES MOTS

Nous pouvons dire qu'un montage électrique ou électronique commence par l'association d'au moins un générateur et d'une charge. Voyons dans cette fiche le comportement correspondant à un réseau constitué d'un générateur et de deux composants passifs identiques ou non (résistances, inductances, condensateurs).

2. CE QU'IL FAUT RETENIR

a) Montages série

L'association en série de deux impédances quelconques correspond au schéma suivant :



Considérons que les deux impédances du circuit proposé sont des résistances :

$$Z_1 = R_1 \text{ et } Z_2 = R_2$$

En appliquant la loi d'Ohm aux bornes de R_1 et de R_2 , nous obtenons les relations suivantes :

$$\begin{cases} V_1 = R_1 \cdot I \\ V_2 = R_2 \cdot I \end{cases} \Rightarrow V = V_1 + V_2 = R_1 \cdot I + R_2 \cdot I = (R_1 + R_2) \cdot I = R_S \cdot I$$

Vu du générateur, la **résistance équivalente** aux deux résistances série R_1 et R_2 est égale à la somme de ces deux résistances. De manière générale, la résistance R_S équivalente à N résistances série est égale à la **somme** des valeurs des résistances placées en série :

$$R_S = \sum_{i=1}^N R_i$$

En reprenant la première relation, il est possible de calculer la valeur de la tension V_1 aux bornes de R_1 en fonction de la tension V appliquée et des valeurs de résistance :

$$I = \frac{V}{R_1 + R_2}$$

$$V_1 = \frac{R_1}{R_1 + R_2} \cdot V$$

C'est du fonctionnement déduit de cette relation que vient le nom du montage **diviseur de tension** : la tension aux bornes d'une des deux résistances est proportionnelle à la tension d'entrée. À noter qu'elle est atténuée dans tous les cas car $R_1 < R_1 + R_2$.

$$Z_s = \sum_{i=1}^N Z_i = \sum_{i=1}^N \frac{1}{Y_i}$$

En remplaçant les deux résistances du montage précédent par **deux inductances en série**, appliquons la même méthode qu'avec le diviseur de tension, mais dans le domaine complexe (fiche 13) et déduisons :

$$\begin{cases} \underline{V}_1 = \underline{Z}_{L1} \cdot \underline{I} = j \cdot L_1 \cdot \omega \cdot \underline{I} \\ \underline{V}_2 = \underline{Z}_{L2} \cdot \underline{I} = j \cdot L_2 \cdot \omega \cdot \underline{I} \end{cases} \Rightarrow \underline{V} = \underline{V}_1 + \underline{V}_2 = j \cdot L_1 \cdot \omega \cdot \underline{I} + j \cdot L_2 \cdot \omega \cdot \underline{I} = j \cdot (L_1 + L_2) \cdot \omega \cdot \underline{I}$$

Voyons comment interpréter la mise en série de deux inductances.

$$\underline{Z}_L = j \cdot L_1 \cdot \omega + j \cdot L_2 \cdot \omega = j \cdot (L_1 + L_2) \cdot \omega \Rightarrow L_s = L_1 + L_2$$

Vu du générateur, l'inductance **équivalente** aux deux inductances série L_1 et L_2 est égale à la somme de ces deux inductances. De manière générale, l'inductance L_s équivalente à N inductances série est égale à la **somme** des valeurs des inductances placées en série :

$$L_s = \sum_{i=1}^N L_i$$

En remplaçant les deux impédances du montage par deux condensateurs, nous pouvons déduire :

$$\begin{cases} \underline{V}_1 = \underline{Z}_{C1} \cdot \underline{I} = \frac{1}{j \cdot C_1 \cdot \omega} \cdot \underline{I} \\ \underline{V}_2 = \underline{Z}_{C2} \cdot \underline{I} = \frac{1}{j \cdot C_2 \cdot \omega} \cdot \underline{I} \end{cases} \Rightarrow \underline{V} = \underline{V}_1 + \underline{V}_2 = \left(\frac{1}{j \cdot C_1 \cdot \omega} + \frac{1}{j \cdot C_2 \cdot \omega} \right) \cdot \underline{I} = \underline{Z}_{C_s} \cdot \underline{I}$$

L'impédance équivalente à la **mise en série de deux condensateurs** n'est pas, à la différence des résistances et des capacités, la somme des capacités mises en jeu.

$$\underline{Z}_{C_s} = \frac{1}{j \cdot C_1 \cdot \omega} + \frac{1}{j \cdot C_2 \cdot \omega} = \frac{C_1 + C_2}{C_1 \cdot C_2 \cdot j \cdot \omega} \Rightarrow C_s = \frac{C_1 \cdot C_2}{C_1 + C_2}$$

De façon générale, la mise en série de N condensateurs est équivalente à une capacité C_s :

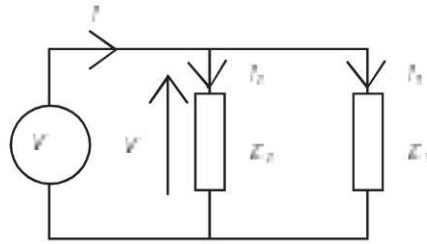
$$C_s = \frac{\prod_{i=1}^N C_i}{\sum_{i=1}^N C_i}$$

Quel que soit le composant retenu (résistance, inductance ou condensateur), la mise en série de N composants est équivalente à une impédance $\underline{Z}_s$:

$$\underline{Z}_P = \sum_{i=1}^N \underline{Z}_i = \sum_{i=1}^N \frac{1}{Y_i}$$

b) Montages parallèle

L'association en parallèle de deux impédances quelconques correspond au schéma suivant :



Considérons que les deux impédances du circuit proposé sont des résistances :

$$Z_1 = R_1 \text{ et } Z_2 = R_2$$

En appliquant la loi d'Ohm aux bornes de R_1 et de R_2 , nous obtenons les relations suivantes :

$$V = I_1 \cdot R_1 = I_2 \cdot R_2 \Rightarrow \begin{cases} I_1 = V/R_1 \\ I_2 = V/R_2 \end{cases}$$

La relation tension-courant est déduite grâce à la loi des nœuds :

$$I = I_1 + I_2 = \frac{V}{R_1} + \frac{V}{R_2} \Rightarrow V = \frac{R_1 \cdot R_2}{R_1 + R_2} \cdot I$$

La résistance équivalente résultant de la mise en parallèle de N résistances peut s'écrire de la

manière suivante en toute généralité : $R_p = \frac{\prod_{i=1}^N R_i}{\sum_{i=1}^N R_i}$

Dans le cas particulier où N résistances mises en parallèle sont identiques de valeur R , la résistance équivalente est égale à R/N .

Il est possible de calculer la valeur du courant aux bornes de R_1 en fonction de la tension appliquée et des valeurs de résistance :

$$I_1 = \frac{V}{R_1} = \frac{\frac{R_1 \cdot R_2}{R_1 + R_2} \cdot V}{R_1} = \frac{R_2}{R_1 + R_2} \cdot V$$

Le montage de mise en parallèle de deux résistances est appelée **diviseur de courant**.

$$Y_p = \sum_{i=1}^N Y_i = \sum_{i=1}^N \frac{1}{Z_i}$$

Remplaçons les deux résistances en parallèle du cas précédent par des inductances :

$$\begin{cases} \underline{I}_1 = \underline{Y}_{L_1} \cdot \underline{V} = \frac{1}{j \cdot L_1 \cdot \omega} \cdot \underline{V} \\ \underline{I}_2 = \underline{Y}_{L_2} \cdot \underline{V} = \frac{1}{j \cdot L_2 \cdot \omega} \cdot \underline{V} \end{cases} \Rightarrow \underline{I} = \underline{I}_1 + \underline{I}_2 = \left(\frac{1}{j \cdot L_1 \cdot \omega} + \frac{1}{j \cdot L_2 \cdot \omega} \right) \cdot \underline{V}$$

$$Y_{Lp} = \frac{1}{j \cdot L_1 \cdot \omega} + \frac{1}{j \cdot L_2 \cdot \omega} = \frac{L_1 + L_2}{L_1 \cdot L_2 \cdot j \cdot \omega} \Rightarrow Z_{Lp} = j \cdot \omega \cdot \frac{L_1 \cdot L_2}{L_1 + L_2} \Rightarrow L_p = \frac{L_1 \cdot L_2}{L_1 + L_2}$$

L'inductance équivalente résultant de la mise en parallèle de N inductances peut s'écrire de la manière suivante en toute généralité :

$$L_P = \frac{\prod_{i=1}^N L_i}{\sum_{i=1}^N L_i}$$

Pour le cas de **deux condensateurs placés en parallèle**, appliquons une méthode de calcul basée sur l'admittance des condensateurs C_1 et C_2 .

$$\begin{cases} \underline{I}_1 = \underline{Y}_C \cdot \underline{V} = j \cdot C_1 \cdot \omega \cdot \underline{V} \\ \underline{I}_2 = \underline{Y}_{C2} \cdot \underline{V} = j \cdot C_2 \cdot \omega \cdot \underline{V} \end{cases} \Rightarrow \underline{I} = \underline{I}_1 + \underline{I}_2 = j \cdot C_1 \cdot \omega \cdot \underline{V} + j \cdot C_2 \cdot \omega \cdot \underline{V} = j \cdot (C_1 + C_2) \cdot \omega \cdot \underline{V}$$

Identifions la capacité équivalente à ces deux condensateurs en parallèle :

$$\underline{Y}_{c_p} = j \cdot (C_1 + C_2) \cdot \omega \Rightarrow C_p = C_1 + C_2$$

En raisonnant en termes d'impédances, nous aurions trouvé :

$$\underline{Z}_{c_p} = \frac{1}{j \cdot (C_1 + C_2) \cdot \omega}$$

La mise en parallèle de N condensateur placés en parallèle est donc la somme des capacités :

$$C_P = \sum_{i=1}^N C_i$$

Ce tableau résume les conclusions de l'étude de cette fiche pour deux composants mis en série ou en parallèle.

	Résistances	Condensateurs	Inductances
Série	$R_1 + R_2$	$\frac{C_1 \cdot C_2}{C_1 + C_2}$	$L_1 + L_2$
Parallèle	$\frac{R_1 \cdot R_2}{R_1 + R_2}$	$C_1 + C_2$	$\frac{L_1 \cdot L_2}{L_1 + L_2}$

3. EN PRATIQUE

En parallèle ou en série, l'association de deux composants de la même famille sert avant tout à obtenir des valeurs de composant non disponibles parmi les valeurs normalisées (fiche 6).

Les calculs relatifs à ces montages sont indispensables pour les théorèmes de Thévenin et de Norton (fiche 24).

L'étude correspondant à deux composants différents dans le montage parallèle ou le montage série fait l'objet de la fiche spécifique sur les filtres (fiche 26).

Le montage diviseur de tension est utilisé sous une forme particulière dans les convertisseurs analogique-numérique de type flash : un CAN flash de N bits comporte 2^N résistances identiques en série (fiche 63).

24 Théorèmes de Thévenin – Norton

Mots clés

Générateur équivalent, impédance équivalente, lois de Kirchhoff.

1. EN QUELQUES MOTS

Ces théorèmes permettent dans certains cas de simplifier une partie du réseau électrique étudié.

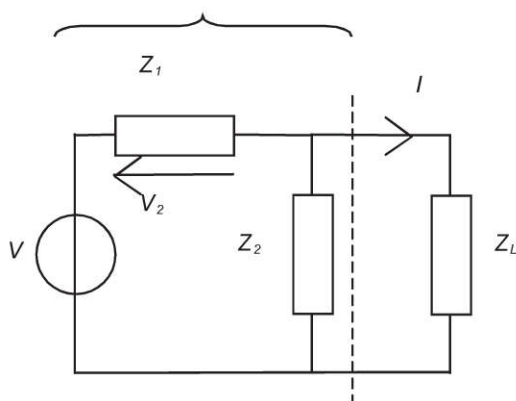
2. CE QU'IL FAUT RETENIR

a) Théorème de Thévenin

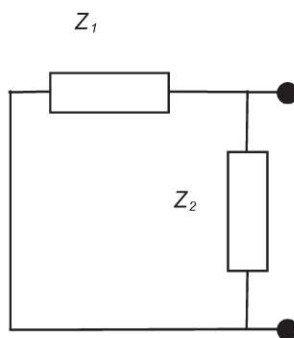
Un réseau, quel que soit le nombre de nœuds, de mailles ou de composants qu'il comporte, peut être vu de deux points A et B distincts et considéré comme étant équivalent à l'association suivante.

- Un générateur de tension dont la force électromotrice est égale à la différence de potentiel entre A et B lorsque le circuit est ouvert entre les points A et B.
- Une impédance série équivalente à celle calculée entre A et B lorsque tous les générateurs et les récepteurs du réseau sont remplacés par leur impédance interne. Chaque générateur de tension idéal est alors remplacé par un court-circuit, chaque générateur de courant idéal est remplacé par un circuit ouvert.

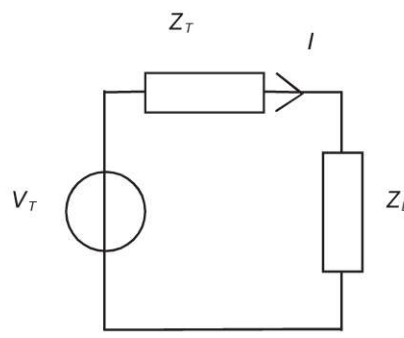
Partie du circuit sur laquelle appliquer le théorème de Thévenin



Calcul de l'impédance équivalente



Générateur de Thévenin équivalent



Pour trouver la tension entre les points A et B lorsque l'impédance de charge Z_L est débranchée, il suffit dans notre cas d'appliquer la formule du pont diviseur de tension (fiche 23) :

$$V_T = V \cdot \frac{Z_2}{Z_1 + Z_2}$$

La deuxième étape consiste à trouver l'impédance équivalente en remplaçant le générateur de tension idéal par un court-circuit. Cette impédance est donc équivalente à la mise en parallèle de Z_1 et de Z_2 :

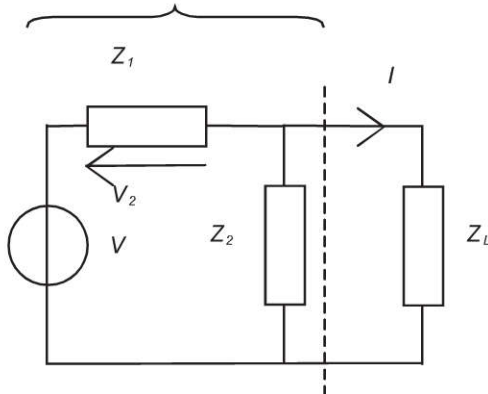
$$Z_T = \frac{Z_1 \cdot Z_2}{Z_1 + Z_2}$$

b) Théorème de Norton

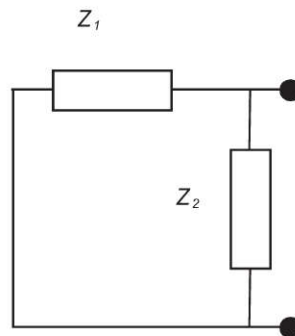
Un réseau, quel que soit le nombre de nœuds, de mailles ou de composants qu'il comporte, peut être vu de deux points A et B distincts et considéré comme étant équivalent à l'association suivante.

- Un générateur de courant égal au courant circulant dans la branche contenant les points A et B lorsqu'un court circuit est placé entre A et B.
- Une impédance parallèle équivalente à celle calculée entre A et B lorsque tous les générateurs et les récepteurs du réseau sont remplacés par leur impédance interne. De même que pour le théorème de Thévenin, chaque générateur de tension idéal est alors remplacé par un court-circuit, chaque générateur de courant idéal est remplacé par un circuit ouvert.

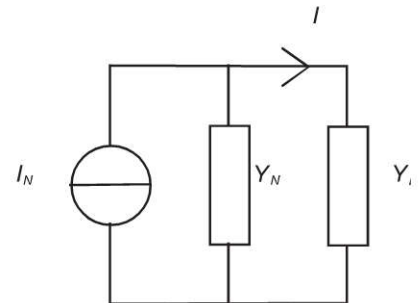
Partie du circuit sur laquelle appliquer le théorème de Norton



Calcul de l'impédance équivalente



Générateur de Norton équivalent



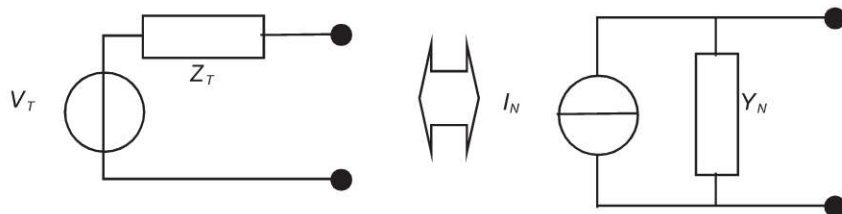
$$I_N = \frac{V}{Z_1}$$

$$Y_N = Y_1 + Y_2 = \frac{1}{Z_1} + \frac{1}{Z_2} \text{ avec } Y_L = \frac{1}{Z_L}$$

3. EN PRATIQUE

Les schémas équivalents de Thévenin et de Norton ne permettent d'effectuer les calculs qu'à l'extérieur des générateurs équivalents et donc à l'extérieur du système.

Il est possible de convertir un circuit de Thévenin en un circuit de Norton et inversement.



On passe directement d'un circuit de Thévenin (figure de gauche) à un circuit de Norton (figure de droite) et inversement, à l'aide des formules suivantes :

Passage de...	Résistance équivalente	Générateur équivalent
Thévenin à Norton	$Z_N = Z_T$	$I_N = \frac{V_T}{Z_T}$
Norton à Thévenin	$Z_T = Z_N$	$V_T = I_N \cdot Z_N$

25 Ponts de mesure

Mots clés

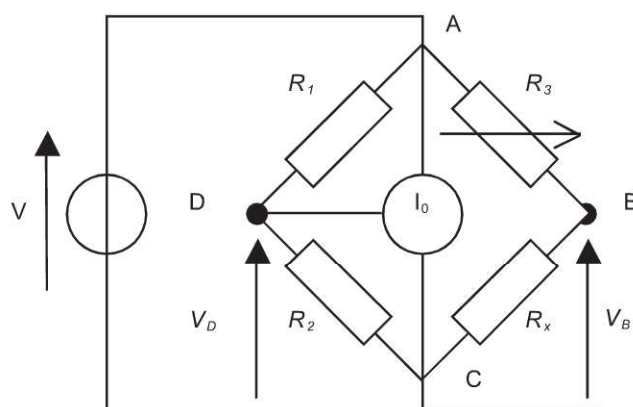
Pont de mesure, galvanomètre, théorème de Millman, ponts de Wheatstone, de Maxwell, de Sauty.

1. EN QUELQUES MOTS

Les ponts de mesure permettent d'évaluer la valeur d'un composant électrique inconnu.

2. CE QU'IL FAUT RETENIR

► Le pont de Wheatstone correspond au dispositif suivant :



La manipulation consiste à équilibrer les deux branches du circuit en pont, avec R_x le composant inconnu, R_1 et R_2 deux résistances connues et R_3 une résistance variable (potentiomètre). Le voltmètre situé entre les nœuds B et D sert à valider l'équilibrage du pont de Wheatstone. Le **théorème de Millman** (fiche 10) permet de calculer la tension V_D :

$$V_D = V \cdot \frac{R_2}{R_1 + R_2}$$

Par analogie, nous pouvons déduire l'expression du potentiel électrique au point B :

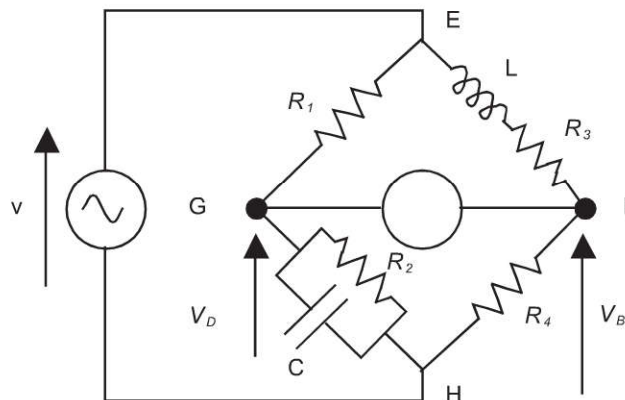
$$V_B = V \cdot \frac{R_x}{R_3 + R_x}$$

Lorsque le pont est équilibré, aucun courant ne passe dans la branche où est situé le galvanomètre. Il n'existe donc plus de différence de potentiel entre les points B et D. En considérant $V_B = V_D$, on trouve :

$$\frac{R_x}{R_3 + R_x} = \frac{R_2}{R_1 + R_2} \Rightarrow R_x = \frac{R_2 \cdot R_3}{R_1}$$

Le principe est le même pour **mesurer des condensateurs** (pont de Sauty) ou **des inductances** (pont de Maxwell). Il faut pour cela que la source de tension continue soit remplacée par une source alternative.

► Le pont de Maxwell se présente sous la forme suivante :



Appliquons le théorème de Millman lorsque le pont est équilibré :

$$R_3 + j \cdot L \cdot \omega = R_1 \cdot R_4 \cdot \left(\frac{1}{R_2} + j \cdot C \cdot \omega \right)$$

En égalisant les parties réelles et imaginaires, nous obtenons des résultats indépendants de la pulsation du générateur :

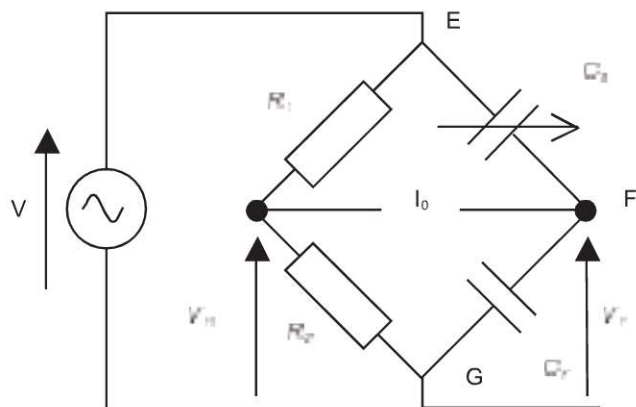
$$R_3 = \frac{R_1 \cdot R_4}{R_2} \text{ et } L = R_1 \cdot R_4 \cdot C$$

► Le pont de Sauty nécessite de remplacer deux des résistances du pont de Wheatstone par des capacités : l'une est variable, l'autre est inconnue.

Similairement au pont de Wheatstone, le théorème de Millman permet de calculer les tensions V_F et V_H :

$$V_H = V \cdot \frac{R_2}{R_1 + R_2} \text{ et}$$

$$V_F = V \cdot \frac{-\frac{j}{C_3 \cdot \omega}}{-\frac{j}{C_X \cdot \omega} - \frac{j}{C_3 \cdot \omega}} = V \cdot \frac{C_3}{C_X + C_3}$$



Considérons $V_F = V_H$, lorsque le pont est équilibré. Nous trouvons alors :

$$\frac{C_3}{C_3 + C_X} = \frac{R_2}{R_1 + R_2} \Rightarrow C_X = \frac{C_3 \cdot R_1}{R_2}$$

3. EN PRATIQUE

Le pont de Wheatstone est souvent utilisé en tant que dispositif de mesure de déformations. Certains matériaux présentent une résistance en fonction de contraintes mécaniques (pression, déformation). Elle permet de fabriquer des capteurs de pression, accélération, etc. La variation de cette résistance est trop faible pour être directement mesurée, le pont de Wheatstone permet d'effectuer ces mesures avec précision.

26 Filtres passifs du premier ordre

Mots clés

Filtre passe-haut, filtre passe-bas, diagramme de Bode, transmittance, atténuation.

1. EN QUELQUES MOTS

L'association de deux composants passifs de même nature a été étudiée dans la [fiche 23](#). Les filtres passifs du premier ordre consistent, eux, à associer deux composants passifs de natures différentes. En toute généralité, un filtre électronique consiste à atténuer un signal en fonction de sa fréquence.

2. CE QU'IL FAUT RETENIR

a) Résolution mathématique dans le domaine réel

Un circuit du premier ordre se modélise sous la forme d'une **équation différentielle d'ordre 1**. Nous sortons alors du cadre des circuits linéaires.

Ces équations différentielles sont de la forme :

$$a_0 \cdot x(t) + a_1 \cdot \frac{dx(t)}{dt} = b$$

L'inconnue est x et a_0 , a_1 et b sont des constantes.

On cherche donc $x(t)$ en la décomposant en l'addition de deux termes :

$$x(t) = x_1(t) + x_2(t)$$

- $x_1(t)$ est la solution générale de **l'équation sans second membre** (parfois notée *essm*) telle que :

$$a_0 \cdot x_1(t) + a_1 \cdot \frac{dx_1(t)}{dt} = 0$$

La solution de cette équation est de la forme :

$$x_1(t) = \pm K \cdot e^{-\frac{t}{\tau}}$$

- $x_2(t)$ est la solution particulière de **l'équation générale avec second membre** (*easm*).

$$a_0 \cdot x(t) = b \Rightarrow x(t) = \frac{a_0}{b} \text{ car } \frac{dx(t)}{dt} = 0 \text{ en régime établi.}$$

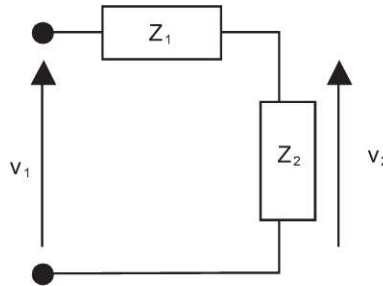
Exemple : soit un circuit RL série (voir 4e figure du tableau suivant)

$$v_1(t) = R \cdot i(t) + L \frac{di(t)}{dt} \text{ soit } v_1(0) = E, \text{ on obtient}$$

$$(\text{essm}) \ v_1(t) = K e^{-\frac{t}{\tau}} \text{ et } v_2(t) = \frac{E}{R} (\text{easm})$$

b) Méthode complexe symbolique

La méthode de résolution complexe, nous allons le voir, nécessite des calculs moins lourds que dans le domaine réel. Cependant, celle-ci ne s'applique que dans le cas et seulement dans le cas du **régime permanent sinusoïdal** et ne tient pas compte du régime transitoire.



Pour calculer T à partir de ce type de schéma, il suffit d'appliquer la formule relative au pont diviseur de tension et de la généraliser aux impédances complexes :

$$v_2 = \frac{Z_2}{Z_1 + Z_2} \cdot v_1 \Rightarrow T = \frac{Z_2}{Z_1 + Z_2}$$

La **transmittance** est un nombre complexe qui relie la sortie à l'entrée d'un circuit (en tension ou en courant).

$$u_2 = T \cdot u_1 = |T| \cdot u_1 \cdot e^{j \cdot \varphi}$$

On considère donc qu'à une fréquence donnée, un signal sinusoïdal est atténué et déphasé par un circuit passif.

Le diagramme de Bode (fiche 19) associé à ce type de filtre se calcule en utilisant la formule suivante :

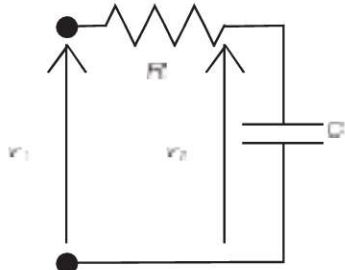
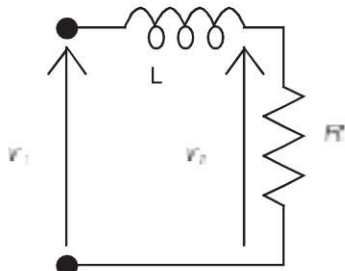
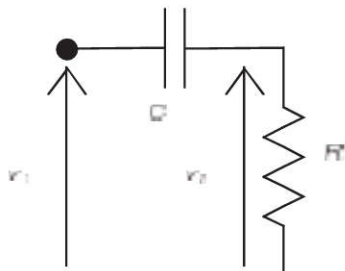
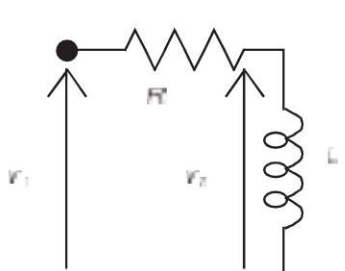
$$T_{dB} = 20 \cdot \log(|T|)$$

Rappelons que l'axe des fréquences utilise une échelle logarithmique afin de représenter aisément plusieurs décades.

Remarques

- Une **décade** est l'intervalle de fréquences situé entre f_0 et $f_1 = 10 \cdot f_0$
- Une **octave** est l'intervalle de fréquences situé entre f_0 et $f_2 = 2 \cdot f_0$

Il s'agit de remplacer les valeurs de Z_1 et de Z_2 par les impédances correspondant aux composants du filtre à étudier. À noter que le cas de l'association d'une inductance et d'un condensateur conduit à l'écriture d'une équation différentielle du second ordre. Ce cas est présenté dans la [fiche 27](#).

Montage	Transmittance complexe	Forme
	$T(\omega) = \frac{1/j \cdot C \cdot \omega}{R + 1/j \cdot C \cdot \omega}$	Filtre passe-bas $T(x) = \frac{1}{1 + j \cdot x}$
	$T(\omega) = \frac{R}{R + j \cdot L \cdot \omega}$	
	$T(\omega) = \frac{R}{R + 1/j \cdot C \cdot \omega}$	Filtre passe-haut $T(x) = \frac{j \cdot x}{1 + j \cdot x}$
	$T(\omega) = \frac{j \cdot L \cdot \omega}{R + j \cdot L \cdot \omega}$	

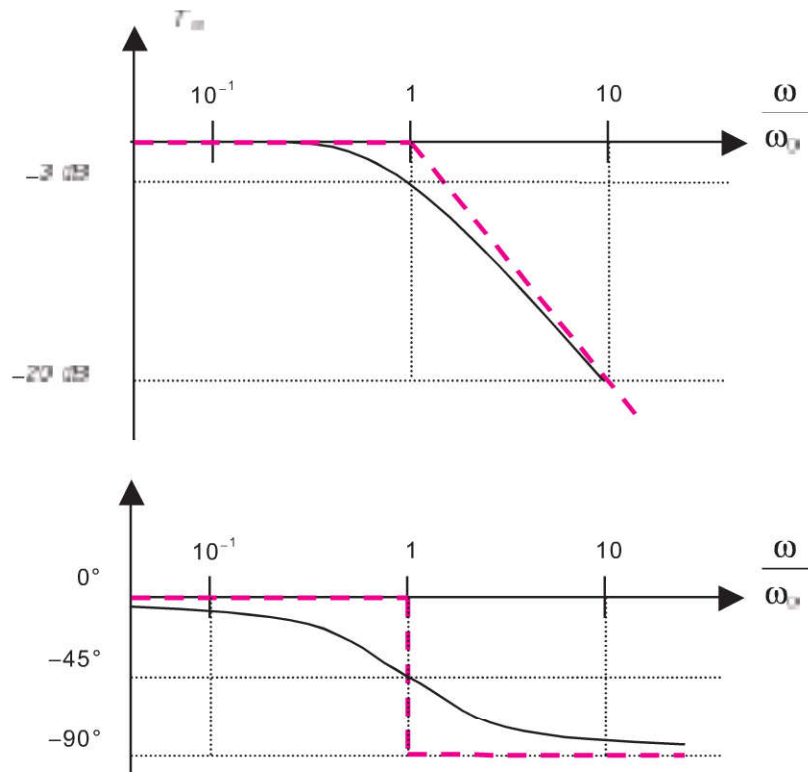
3. EN PRATIQUE

Pour un circuit RC série avec, tension prise aux bornes du condensateur, la transmittance complexe s'écrit :

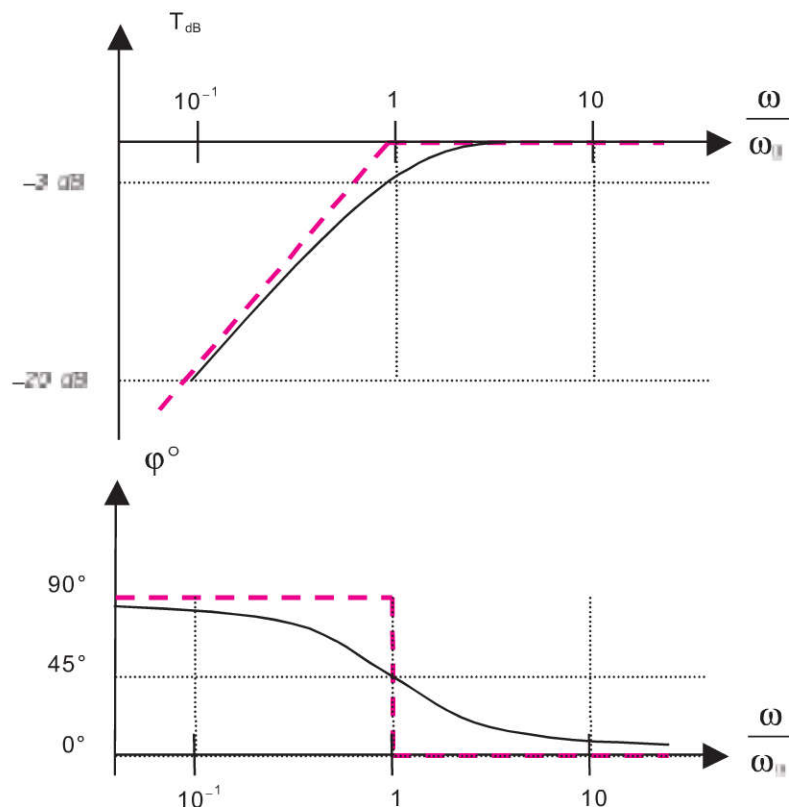
$$T(\omega) = \frac{1/j \cdot C \cdot \omega}{R + 1/j \cdot C \cdot \omega} = \frac{1}{1 + j \cdot R \cdot C \cdot \omega}$$

L'étude analytique du module de la transmittance montre que les fréquences atténuées de moins de 3 dB sont en-deçà de la fréquence de coupure, et lorsque le signal d'entrée a une fréquence au-delà de la fréquence de coupure, l'atténuation de filtre augmente. Pour cette raison,

le filtre considéré est dit passe-bas. L'étude complète de ce filtre est présentée sur la [fiche 19](#). Le **diagramme de Bode** correspondant à ce type de filtre est présenté dans les diagrammes **semi-logarithmiques** suivants.



L'étude d'un filtre passe-haut du premier ordre (filtre RC sortie sur la résistance ou filtre RL sortie sur l'inductance) aboutirait à des conclusions réciproques : ce sont les fréquences les plus hautes qui sont les moins atténuées. Le diagramme de Bode est alors le suivant :



27 Filtres passifs du second ordre

Mots clés

Circuits RLC série, filtre, facteur de qualité, pulsation réduite, fréquence de résonance, transmittance.

1. EN QUELQUES MOTS

Un circuit passif du second ordre se modélise sous la forme d'une **équation différentielle d'ordre 2**. Pour parvenir à ce résultat, il faut utiliser dans un même circuit les trois composants passifs étudiés dans la **fiche 5**, c'est-à-dire l'association en parallèle ou en série d'une ou plusieurs résistances R , condensateurs C et inductances L . C'est pourquoi ces filtres sont généralement appelés filtres RLC. Dans cette fiche, nous étudions les circuits RLC série en **régime sinusoïdal forcé**.

2. CE QU'IL FAUT RETENIR

Soient $\underline{Z}_1$, $\underline{Z}_2$ et $\underline{Z}_3$ les impédances complexes de trois composants passifs de types différents montés en série. Le générateur est une source de tension sinusoïdale v_1 .

L'impédance $\underline{Z}$ du montage vue de v_1 (noté $\underline{U}$ dans le domaine complexe) est la somme des impédances (montage série) :

$$\underline{Z} = \underline{Z}_1 + \underline{Z}_2 + \underline{Z}_3$$

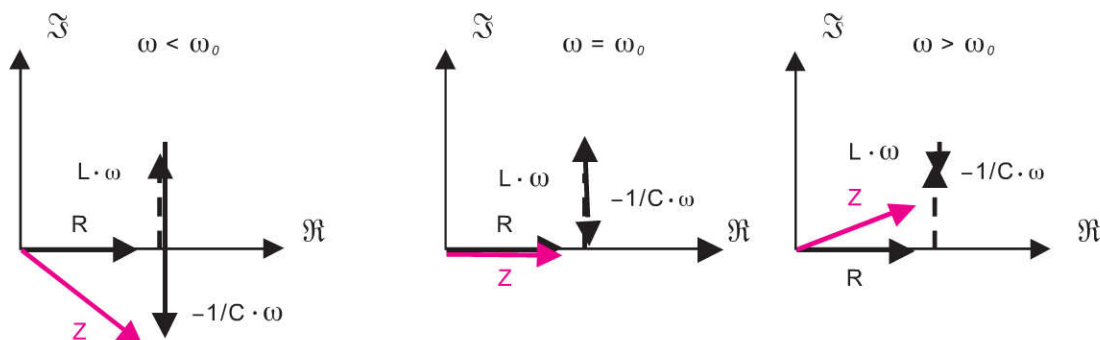
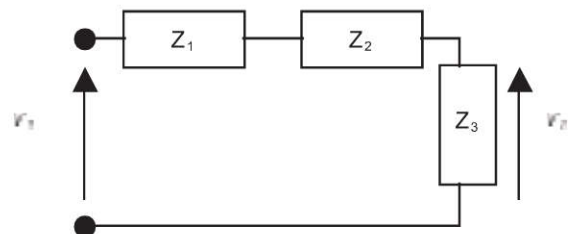
Pour calculer la transmittance $T(\omega)$ à partir de ce type de schéma, il suffit d'appliquer la formule relative au pont diviseur de tension et de la généraliser aux impédances complexes :

$$v_2 = \frac{\underline{Z}_3}{\underline{Z}_1 + \underline{Z}_2 + \underline{Z}_3} \cdot v_1 \Rightarrow T = \frac{\underline{Z}_3}{\underline{Z}_1 + \underline{Z}_2 + \underline{Z}_3}$$

Soit $\underline{Z}_1 = j \cdot L \cdot \omega$, $\underline{Z}_2 = R$ et $\underline{Z}_3 = \frac{-j}{C \cdot \omega}$, on obtient

$$\underline{Z} = R + j \cdot \left(L \cdot \omega - \frac{1}{C \cdot \omega} \right), \quad Z = |\underline{Z}| = \sqrt{R^2 + \left(L \cdot \omega - \frac{1}{C \cdot \omega} \right)^2}$$

Représentons l'impédance complexe du système RLC série sous forme de **diagramme de Fresnel**.



Avant la résonance, à $\omega < \omega_0$, l'impédance (purement imaginaire) du condensateur est plus grande, en valeur absolue, que celle de l'inductance. Après la résonance, à $\omega > \omega_0$, l'impédance de l'inductance est plus grande, en valeur absolue, que celle du condensateur. À la résonance, ces impédances se compensent : l'impédance globale du système est purement résistive.

La fréquence de résonance ($f_0 = \omega_0 / 2 \cdot \pi$) intervient lorsque la partie imaginaire de $\underline{Z}$ est nulle : les composantes inductive et capacitive se compensent. Le module de l'impédance est alors minimum.

$$L \cdot \omega_0 = \frac{1}{C \cdot \omega_0} \Rightarrow L \cdot C \cdot \omega_0^2 = 1 \Rightarrow \omega_0 = \frac{1}{\sqrt{L \cdot C}} \Rightarrow f_0 = \frac{1}{2 \cdot \pi \cdot \sqrt{L \cdot C}}$$

Dans ce cas, le courant circulant dans le réseau est maximum.

$$\underline{I} = \frac{\underline{U}}{R + j \cdot \left(L \cdot \omega - \frac{1}{C \cdot \omega} \right)} = \frac{\underline{U}}{R + j \cdot Q \cdot R \cdot \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)}$$

Q est le facteur de qualité du système : $Q = \frac{L \cdot \omega_0}{R} = \frac{1}{R \cdot C \cdot \omega_0}$

En introduisant la pulsation réduite $x = \omega / \omega_0$, le module et la phase de $\underline{I}$ s'écrivent :

$$|\underline{I}| = \frac{|\underline{U}|}{R \cdot \sqrt{1 + Q^2 \left(x - \frac{1}{x} \right)^2}} \quad \text{et} \quad \varphi = \text{Arc tan} \left[-Q \cdot \left(x - \frac{1}{x} \right) \right]$$

À la pulsation de résonance, intensité et courant sont en phase, l'impédance équivalente étant purement résistive. Avant la résonance, le courant est en avance de phase sur la tension car l'impédance est plus capacitive qu'inductive et réciproquement après la résonance.

La bande passante $\Delta\omega$ correspond à la largeur de bande pour laquelle l'intensité est supérieure à $|\underline{I}|_{\max} / \sqrt{2}$.

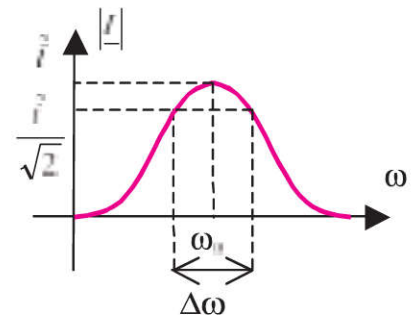
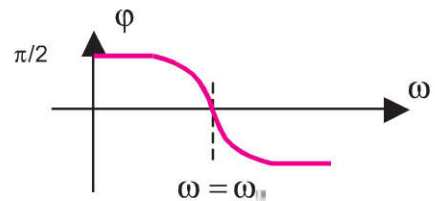
$$\Delta\omega = \frac{R}{L} = \frac{\omega_0}{Q} \Leftrightarrow \Delta f = \frac{1}{2 \cdot \pi} \cdot \frac{R}{L} = \frac{f_0}{Q}$$

Afin d'étudier la résonance en tension, calculons $\underline{U}_C$ aux bornes du condensateur :

$$\underline{U}_C = -\frac{j}{C \cdot \omega} \cdot \underline{I} \Rightarrow \underline{U}_C = \frac{\underline{U}}{1 - L \cdot C \cdot \omega^2 + j \cdot R \cdot C \cdot \omega}$$

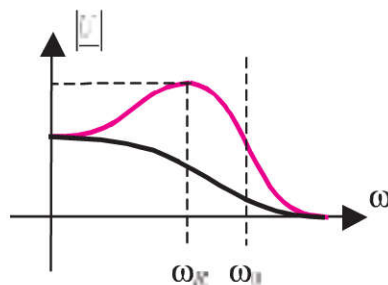
En utilisant la pulsation réduite x et le facteur de qualité Q , nous trouvons :

$$\underline{U}_C = \frac{\underline{U}}{1 - x^2 + j \cdot \frac{x}{Q}} \Rightarrow |\underline{U}_C| = \frac{|\underline{U}|}{\sqrt{(1 - x^2)^2 + \left(\frac{x}{Q} \right)^2}}$$



Étudions le comportement de $|\underline{U}_C|$, et en particulier la pulsation de résonance ω_R en fonction du coefficient de qualité Q :

- $Q < 1/\sqrt{2} \Rightarrow \omega_R = 0$, il n'y a donc pas de résonance en tension (courbe noire)
- $Q > 1/\sqrt{2} \Rightarrow \omega_R = \omega_0 \cdot \sqrt{1 - \frac{1}{Q^2}}$ (courbe rouge). Dans le cas particulier où $Q > 10$, on obtient $\omega_R \approx \omega_0$



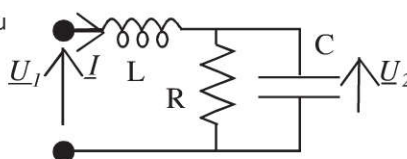
3. EN PRATIQUE

La tension peut également être prélevée aux bornes de la résistance ou de l'inductance, la transmittance complexe est alors modifiée et par là même la forme du filtre.

Montage	Transmittance complexe	Forme
	$T(x) = \frac{1}{1 - x^2 + j \cdot \frac{x}{Q}}$	Filtre passe-bas
	$T(x) = \frac{-x^2}{1 - x^2 + j \cdot \frac{x}{Q}}$	Filtre passe-haut
	$T(x) = \frac{-j \cdot \frac{x}{Q}}{1 - x^2 + j \cdot \frac{x}{Q}}$	Filtre passe-bande

Étudions le circuit RLC suivant et traçons son diagramme de Bode (fiche 19) :

$R = 30 \, \Omega$ ou $300 \, \Omega$ ou
 $3 \, \text{k}\Omega$
 $L = 1 \, \text{mH}$
 $C = 1 \, \mu\text{F}$



L'impédance totale de ce filtre vue de $\underline{U}_1$ est équivalente à la mise en série de l'inductance avec un filtre RC parallèle :

$$\underline{Z} = \frac{\underline{U}_1}{\underline{I}} = \underline{Z}_L + R // \underline{Z}_C = j \cdot L \cdot \omega + \frac{R}{1 + j \cdot R \cdot C \cdot \omega} \quad \text{car } R // \underline{Z}_C = \frac{R \cdot \frac{1}{jC\omega}}{R + \frac{1}{jC\omega}} = \frac{R}{1 + jRC\omega}$$

La fonction de transfert de ce filtre est la suivante :

$$T(\omega) = \frac{1}{1 - L \cdot C \cdot \omega^2 + j \cdot \frac{L \cdot \omega}{R}} \Rightarrow T(x) = \frac{1}{1 - x^2 + \frac{j}{Q} \cdot x} \text{ avec } Q = \frac{R}{L \cdot \omega_0} \text{ et } \omega_0 = \frac{1}{\sqrt{L \cdot C}}$$

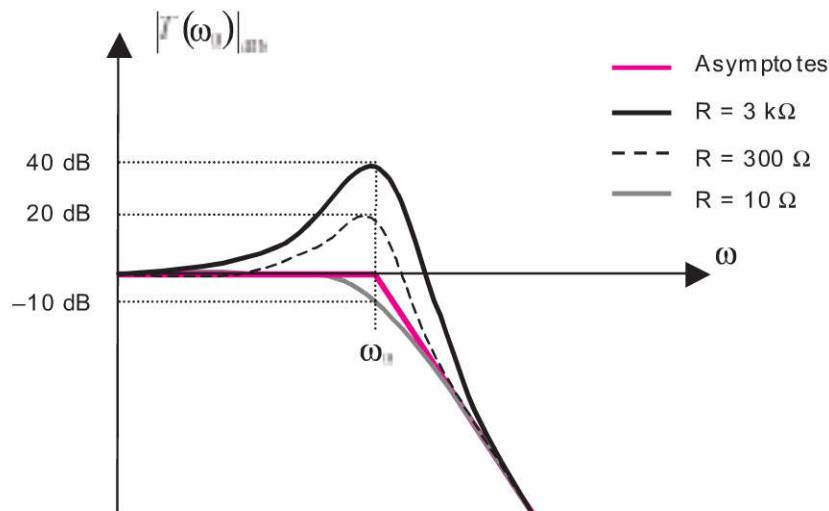
En module, la fonction de transfert devient $|T(\omega)| = \frac{1}{\sqrt{\left(1 - \left(\frac{\omega}{\omega_0}\right)^2\right)^2 + \left(\frac{\omega}{Q \cdot \omega_0}\right)^2}}$

Étudions le comportement asymptotique de $|T(\omega)|$:

- pour $\omega \ll \omega_0$, $|T(\omega)| \approx 1 \Rightarrow |T(\omega)|_{\text{dB}} \approx 0$ (C se comporte comme un circuit ouvert $\rightarrow$ circuit RL série)
- pour $\omega = \omega_0$, $|T(\omega)| \approx Q \Rightarrow |T(\omega)|_{\text{dB}} \approx 20 \cdot \log(Q)$
- pour $\omega \gg \omega_0$, $|T(\omega)| \approx (\omega_0/\omega)^2 \Rightarrow |T(\omega)|_{\text{dB}} \approx 40 \cdot \log(\omega_0/\omega)$ (L se comporte comme un circuit ouvert et C comme un court-circuit)

Cette étude montre que le filtre est un passe-bas. Au-delà de la fréquence de coupure, le filtre atténue de 40 dB de plus à $\omega_2 = 10 \cdot \omega_1$ qu'à ω_1 : le module de la fonction de transfert est donc tangent à une équation de droite de pente 40 dB par décade.

R	3 k Ω	300 Ω	30 Ω	10 Ω
Q	95	9,5	0,95	0,32
$ T(\omega_0) _{\text{dB}} \approx 20 \cdot \log(Q)$	40 dB	20 dB	-0,4 dB	-10 dB



28 Matériaux semi-conducteurs

Mots clés

Silicium, Germanium, électron, trou, conduction, résistivité, semi-conducteurs intrinsèques, métaux, isolants

1. EN QUELQUES MOTS

Les matériaux semi-conducteurs sont à la base de la plupart des composants électroniques : diodes, transistors et, par extension, circuits intégrés.

2. CE QU'IL FAUT RETENIR

a) Conductivité

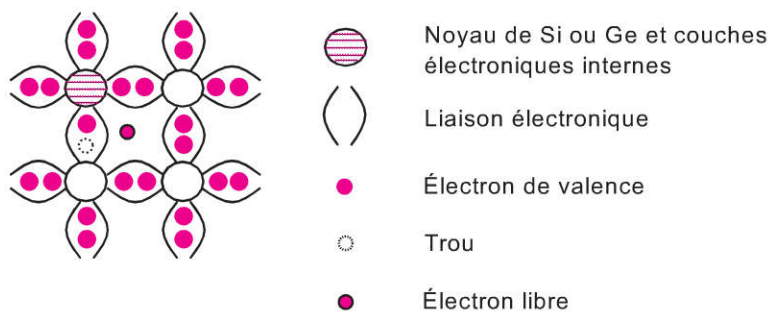
La conductivité (σ en $S \cdot m^{-1}$) est l'inverse de la résistivité (ρ en $\Omega \cdot m$). La conductivité est donc la **capacité d'un matériau à laisser circuler les charges électriques**. Les matériaux peuvent être **classés selon leur conductivité**. On dégage alors trois grandes familles : les **isolants** qui ne laissent pas passer le courant, les **métaux** qui ont un grand nombre d'électrons libres pour la circulation du courant, et les **semi-conducteurs** qui ont un comportement intermédiaire.

$\sigma \text{ (S m}^{-1}\text{)}$	10^6	10^3	10^{-3}	10^{-8}
	10^{-6}	10^{-3}	10^3	10^8
	Métaux	Semiconducteurs	Isolants	
	$\rho \text{ (}\Omega\cdot\text{m)}$			

b) Semi-conducteurs intrinsèques

Les semi-conducteurs sont situés dans la 4^e colonne de la classification périodique des éléments, ce qui signifie qu'ils présentent **quatre électrons de valences** (c'est-à-dire quatre électrons sur la couche électronique la plus élevée). Un semi-conducteur est dit **intrinsèque** lorsqu'il est pur, c'est-à-dire qu'**aucun autre élément chimique** n'est présent dans le cristal considéré.

Couche électronique	1	2	3	4
Si (14 e ⁻)	2	8	4	0
Ge (32 e ⁻)	2	8	18	4



À la température de 0 °K (« zéro absolu », soit environ -273,15 °C), les semi-conducteurs sont des isolants parfaits, ils ne contiennent aucun électron libre. Lorsque la température augmente, des liaisons covalentes se rompent : des électrons de valence quittent leur liaison. Le nombre d'électrons libres augmente exponentiellement avec la température.

La place laissée libre dans une liaison s'appelle un **trou**. Cette vacuité équivaut à une **charge positive** qui, elle aussi, peut se déplacer dans le réseau cristallin. Il y a **autant d'électrons libres que de trous**. Soient :

- n_i la concentration en électrons libres par unité de volume ;
- p_i la concentration en trous par unité de volume.

Dans un semi-conducteur intrinsèque, $n_i = p_i$

Matériau	Cu	Ge	Si	Silice SiO ₂
Famille électrique	Conducteur	Semi-conducteur		Isolant
Nombre atomique	29	32	14	
Masse molaire (g·mol ⁻¹)	63,55	72,64	28,1	60,08
Masse spécifique (g·cm ⁻³)	8,92	5,32	2,33	2,634
Nombre d'atomes par cm ³	8,45·10 ²²	4,42·10 ²²	5·10 ²²	2,35·10 ²²
n_i par cm ³ @ 300 K		2,4·10 ¹³	1,45·10 ¹⁰	
Mobilité des électrons (cm ² ·V ⁻¹ ·s ⁻¹)		3 900	1500	
Mobilité des trous (cm ² ·V ⁻¹ ·s ⁻¹)		1 900	475	
Résistivité @ 300 K (Ω·m)	16,78·10 ⁻⁹	4,7·10 ³	2,3·10 ⁷	10 ¹⁶

Un métal ne présente qu'un seul type de porteurs : les électrons. La conductivité σ ne dépend alors que de la mobilité des électrons :

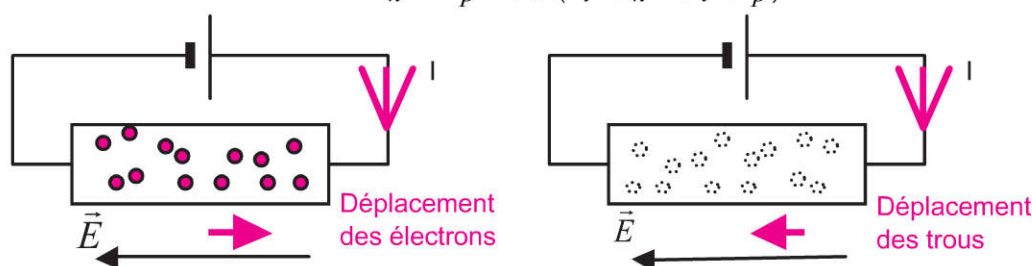
$$\sigma = \frac{1}{\rho} = |e| \cdot n \cdot \mu_n \text{ avec } n \text{ la concentration}$$

Note

Par la suite, dans les différentes notations, l'indice n signifie « négatif » et l'indice p « positif ».

Dans un semi-conducteur, le phénomène de conduction est dû à l'action combinée des électrons et des trous. La conductivité dépend donc de la mobilité μ_n des électrons, mais aussi de celle des trous (μ_p).

$$\sigma = \sigma_n + \sigma_p = |e| \cdot (n_i \cdot \mu_n + p_i \cdot \mu_p)$$



À noter que les concentrations de porteurs de charge n_i et p_i sont identiques. Par contre, la mobilité des électrons est supérieure à la mobilité des trous ($\mu_n > \mu_p$).

3. EN PRATIQUE

Le silicium est le semi-conducteur le plus répandu du fait de son abondance naturelle. Certains matériaux semi-conducteurs tels que l'arséniure de gallium (AsGa) ou le carbure de silicium (CSi) sont constitués de plusieurs éléments chimiques. Afin de fabriquer des composants électroniques, ces matériaux ne sont pas utilisés purs, mais on insère d'autres atomes dans leur structure afin de modifier leur comportement électrique : ce sont les **semi-conducteurs extrinsèques** qui font l'objet de la [fiche 29](#).

29 Semi-conducteurs extrinsèques

Mots clés

Dopage, semi-conducteur type P, semi-conducteur type N.

1. EN QUELQUES MOTS

En électronique, les semi-conducteurs ne sont pas utilisés purs. Afin de modifier leurs propriétés électriques, et principalement la concentration des porteurs libres, des atomes d'espèces chimiques différentes sont introduits dans la structure cristalline des semi-conducteurs intrinsèques.

2. CE QU'IL FAUT RETENIR

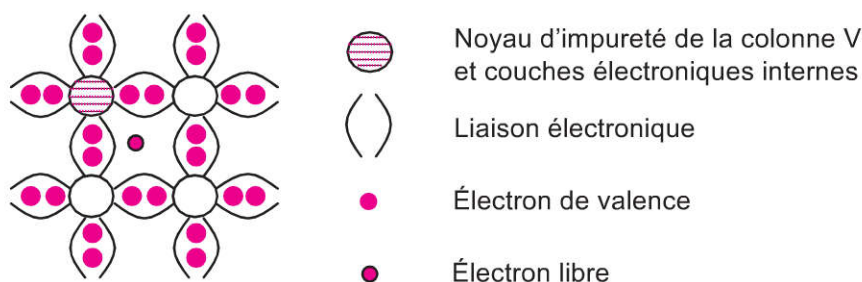
On obtient un semi-conducteur extrinsèque par l'ajout d'impuretés dans le réseau cristallin d'un semi-conducteur intrinsèque. Ces impuretés, appelées **dopants**, appartiennent généralement aux colonnes III et V, voisines de la colonne IV des semi-conducteurs que sont le Si et le Ge.

- **Les éléments de la colonne III** ont trois électrons de valence, soit un de moins que les atomes de la colonne IV. Par conséquent, introduire dans un semi-conducteur pur des éléments de la colonne III tels que du Bore ou du Gallium revient à **rajouter des trous**.
- **Les éléments de la colonne V** ont cinq électrons de valence, soit un de plus que les atomes de la colonne IV. L'introduction d'éléments de la colonne V tels que de l'Arsenic ou du Phosphore dans un semi-conducteur pur revient donc à **rajouter des électrons libres**.

Il est important de noter que **la charge globale d'un semi-conducteur extrinsèque est nulle** : les éléments dopants sont électriquement neutres, la charge des porteurs libres est compensée par la charge des noyaux des impuretés introduites.

a) Semi-conducteurs de type N

Le dopant introduit dans un semi-conducteur de type N appartient à la colonne V. **Chaque atome introduit apporte une charge négative mobile** (un électron) et une **charge positive fixe** qui ne participe pas à la conduction. Soit N_d la concentration en atomes donneurs d'électrons.



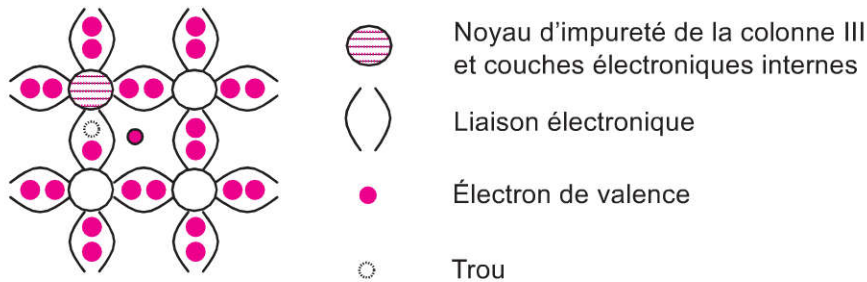
	$\forall N_d$	Si $N_d \gg n_i$
Concentration en charges fixes	N_d	N_d
Concentration n en électrons libres	$N_d + n_i$	$\sim N_d$
Concentration p en trous	p_i	$\ll n$

Dans un semi-conducteur de type N, la concentration en électrons libres est très supérieure à la concentration en trous. En première approximation, la conductivité équivaut donc à celle des électrons libres :

$$\sigma \approx \sigma_n = N_d \cdot |e| \cdot \mu_n$$

b) Semi-conducteurs de type P

Le dopant introduit dans un semi-conducteur de type P appartient à la colonne III. Chaque atome introduit apporte une charge positive mobile (un trou) et une charge négative fixe qui ne participe pas à la conduction. Soit N_a la concentration en atomes accepteurs d'électrons.



	$\forall N_a$	Si $N_a \gg p_i$
Concentration en charges fixes	N_a	N_a
Concentration n en électrons libres	n_i	$\ll p$
Concentration p en trous	$N_a + p_i$	$\sim N_a$

Dans un semi-conducteur de type P, la concentration en trous est très supérieure à la concentration en électrons libres. En première approximation, la conductivité équivaut donc à celle des trous :

$$\sigma \approx \sigma_p = N_a \cdot |e| \cdot \mu_p$$

3. EN PRATIQUE

Produire un semi-conducteur extrinsèque suppose que des impuretés sont introduites dans la structure cristalline. Les deux techniques principales de dopage sont la **diffusion** et l'**implantation ionique**.

a) Diffusion

Le dopage par diffusion consiste à incorporer par **migration** les impuretés dans le semi-conducteur. Les lois de la diffusion, dites **lois de Fick**, régissent le déplacement d'une espèce chimique dans un milieu des plus fortes concentrations vers les plus faibles. Ce déplacement est provoqué par le phénomène d'agitation thermique : plus la température est élevée, plus la diffusion est rapide. Un cristal tel que le silicium n'a jamais une structure parfaite : ce sont les imperfections de la structure cristalline qui permettent à la diffusion d'avoir lieu.

Ce type de dopage se déroule en deux temps. Tout d'abord, une couche très fine de substance dopante est déposée en surface du matériau semi-conducteur. La redistribution est l'étape de diffusion en elle-même : les atomes d'impureté migrent en profondeur dans le semi-conducteur sous haute température (de l'ordre de 1 000 °C).

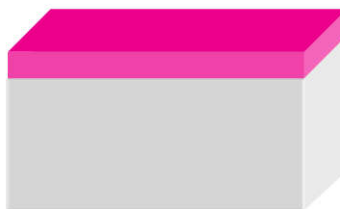
	Conducteurs	SC intrinsèques		SC extrinsèques	
				Type N	Type P
Porteurs Concentration	électrons n	électrons $n_n = n$	trous $n_p = p$	Électrons majoritaires $n \approx N_d$ Trous minoritaires $p = \frac{n_i^2}{N_d}$	Trous majoritaires $p \approx N_a$ Électrons minoritaires $n = \frac{n_i^2}{N_a}$
Mobilité	μ	μ_n	μ_p	μ_n	μ_p
Vitesse d'entraînement	$\vec{v} = -\mu \cdot \vec{E}$	$\vec{v}_n = -\mu_n \cdot \vec{E}$	$\vec{v}_p = \mu_p \cdot \vec{E}$		
Densité de courant	$\vec{j} = -n \cdot e \cdot \vec{v} = \sigma \cdot \vec{E}$	$\vec{j} = -n \cdot e \cdot \vec{v}_n + p \cdot e \cdot \vec{v}_p = \sigma_i \cdot \vec{E}$		$\vec{j} \approx \vec{j}_n = -n \cdot e \cdot \vec{v}_n$	$\vec{j} \approx \vec{j}_p = p \cdot e \cdot \vec{v}_p$
Conductivité	$\sigma = n \cdot e \cdot \mu$	$\sigma_i = n_i \cdot e \cdot (\mu_n + \mu_p)$		$\sigma \approx \sigma_n = N_d \cdot e \cdot \mu_n$	$\sigma \approx \sigma_p = N_a \cdot e \cdot \mu_p$

Remarque

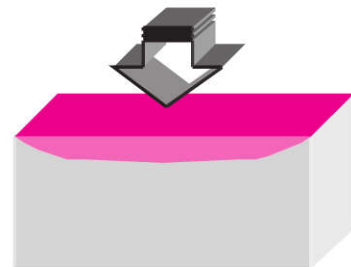
Le cristal reste électriquement neutre : il comporte autant de charges électriques positives que négatives, mais certaines sont fixes et d'autres sont mobiles. Le phénomène de conduction est dû aux électrons et aux trous.



Semi-conducteur intrinsèque



Dépôt d'une couche mince de dopant



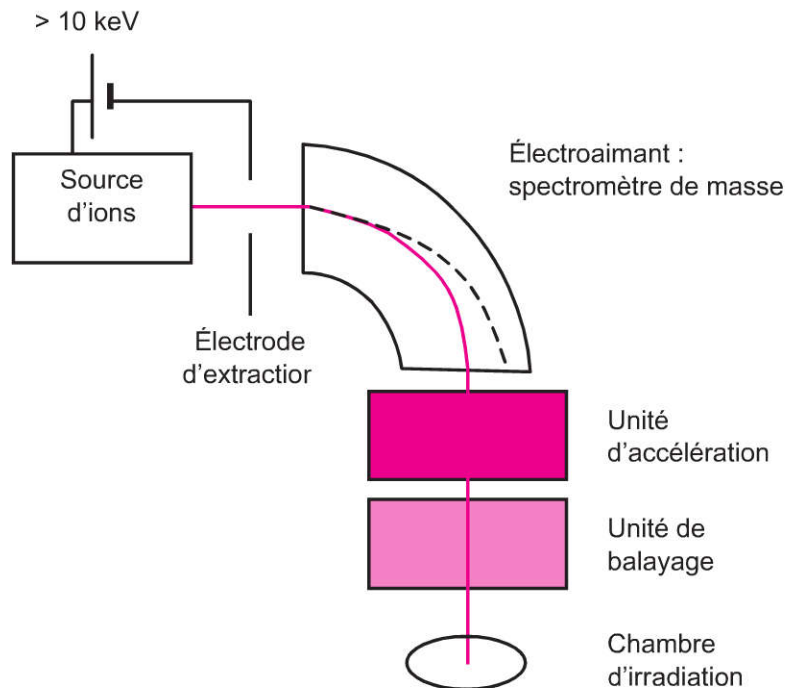
Diffusion du dopant

b) Implantation ionique

Ce procédé consiste à faire entrer des ions dans la matière par accélération. Les limites de solubilité peuvent être dépassées. Une étape de recuit thermique vise à activer le dopage.

Il présente l'avantage de pouvoir incorporer n'importe quel atome dans la matière, avec une très grande précision en termes de quantité et de localisation, mais il a l'inconvénient de nécessiter des installations lourdes et coûteuses.

Les ions sont émis par une source puis triés par un électroaimant : en fonction de la masse des particules et du champ magnétique créé, les ions sont plus ou moins déviés. La maîtrise de cette déviation permet de sélectionner les seuls ions qui seront accélérés et guidés vers la cible qui est la surface du semi-conducteur à doper.



30 Jonction PN et diodes

Mots clés

Champ électrique, semi-conducteurs extrinsèques, polarisations directe et inverse.

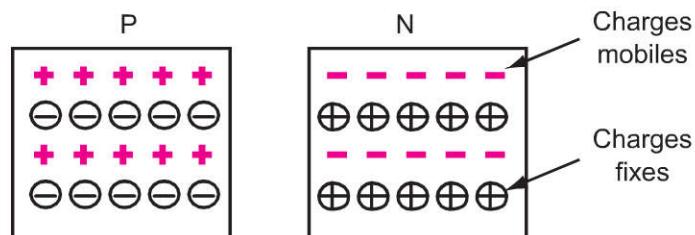
1. EN QUELQUES MOTS

Dans les [fiches 28 et 29](#), nous avons vu comment se comportent les semi-conducteurs intrinsèques et extrinsèques. Ces principes physiques sont à la **base de l'électronique moderne** : l'association de semi-conducteurs de type N et P permet la fabrication des diodes à jonction ainsi que des transistors bipolaires ([fiche 32](#)) et à effet de champ ([fiche 42](#)).

2. CE QU'IL FAUT RETENIR

a) Jonction PN à l'équilibre

Considérons **deux blocs de semi-conducteurs extrinsèques** : le premier de **type P**, le deuxième de **type N**.



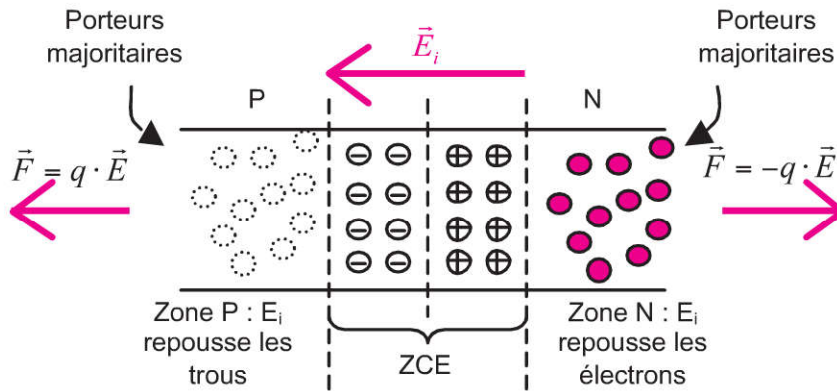
Le bilan des concentrations en porteurs de charge dans chacun des cristaux semi-conducteurs dopés P et N, résumé dans le tableau suivant, est explicité dans la [fiche 29](#).

	SC dopé P	SC dopé N
Porteurs majoritaires	pp # Na	nn # Nd
Porteurs minoritaires	np	pn
Impuretés	Na	Nd
Charges fixes	Négatives Na	Positives Nd

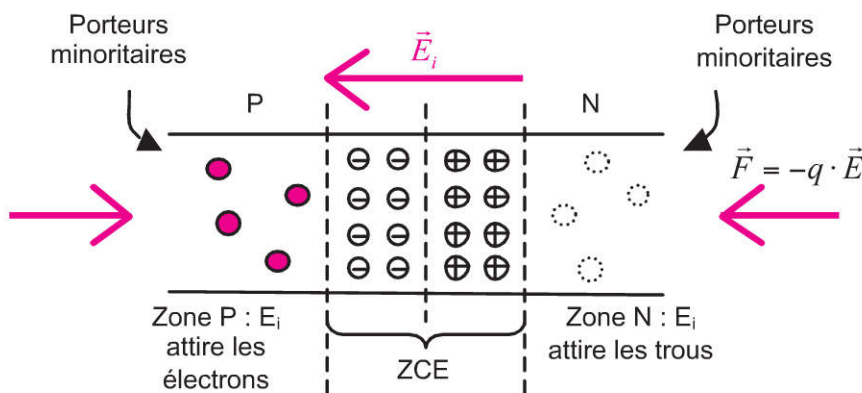
Si les deux cristaux se touchent, la zone de contact est le siège des phénomènes suivants :

- les électrons libres du semi-conducteur de type N se recombinent avec les trous du semi-conducteur de type P. De ce côté de la jonction, la charge devient donc négative ;
- les trous du semi-conducteur de type P se recombinent avec les électrons libres du semi-conducteur de type N. De ce côté de la jonction, la charge devient donc positive.

De part et d'autre de la jonction, on obtient une zone dépourvue de porteurs libres, c'est ce que l'on appelle la **zone désertée** ou encore **zone d'appauvrissement**, **zone de transition**, **zone de charge d'espace** (ZCE). Il y apparaît un champ électrique interne dû aux charges fixes opposées. Dans la zone désertée, les charges négatives sont égales aux charges positives.



Le champ électrique créé à la jonction repousse les porteurs majoritaires. Sans intervention extérieure, ils ne peuvent pas franchir la barrière de potentiel. Ce même champ électrique interne E_i favorise le passage des porteurs minoritaires.



Remarque

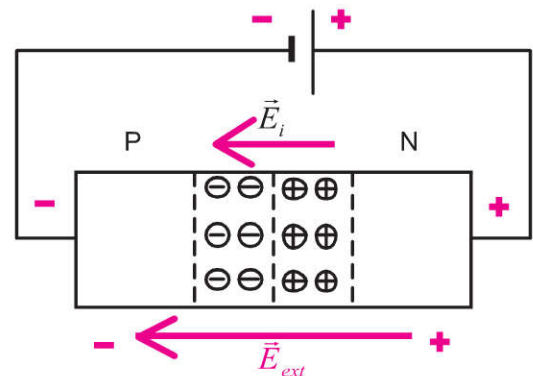
Le très faible courant dû aux porteurs minoritaires est compensé par un courant dû aux porteurs majoritaires. Ces deux courants de sens contraires s'annulent. Dans la zone désertée, les charges négatives sont au même nombre que les charges positives. La zone désertée s'étend vers la zone la plus faiblement dopée.

b) Jonction PN polarisée par une force électromotrice extérieure

Lorsqu'une **force électromotrice est appliquée** aux bornes de la jonction PN, un **champ électrique supplémentaire** se superpose au champ électrique interne à l'équilibre. Le champ électrique total est égal à la somme des champs électrique interne et externe soit : $\vec{E}_{total} = \vec{E}_i + \vec{E}_{ext}$. Selon le sens de la fem extérieure, nous parlerons de **polarisation inverse** ou de **polarisation directe**.

Lorsque le **champ électrique** dû à la fem appliquée est de **même sens que le champ électrique** interne, la **polarisation de la jonction est inverse**.

Dans ce cas, la zone de charge d'espace s'étend. Les porteurs majoritaires ne peuvent plus traverser la barrière, mais tous les porteurs minoritaires la traversent. Le courant est donc uniquement dû aux **porteurs minoritaires**. Ce **courant inverse**, ou de **saturation**, est constant quelle que soit la tension extérieure appliquée, mais il dépend de la température. Cependant, la tension de polarisation extérieure inverse doit rester sous un certain seuil V_Z au-delà duquel la jonction



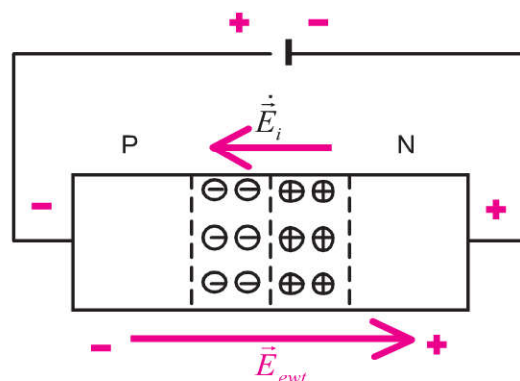
peut « claquer » : le composant est alors détruit. La zone de tension comprise entre V_Z et la tension de claquage est appelée **zone d'avalanche**.

En **polarisation directe**, le champ électrique externe est de sens opposé au champ électrique interne.

Lorsque la tension extérieure augmente, la barrière de tension diminue puis elle disparaît au-delà d'un certain seuil V_γ . Ce courant, appelé **courant de diffusion**, est dû aux **porteurs majoritaires**. V_γ est d'environ 0,3 V pour une diode au germanium et 0,6 V pour une diode au silicium.

Remarque

dans le circuit extérieur, le courant est dû uniquement aux électrons. Lorsqu'un trou parvient au pôle négatif, il se recombine avec un électron qui est absorbé à l'autre extrémité P du cristal.



Polarisation	Porteurs	Variation du courant
Inverse	Courant dû aux porteurs minoritaires	I_s très faible, constant quel que soit $V_Z < V_{PN} < 0$ Dépend de la température
Directe	Courant dû aux porteurs majoritaires	I varie de façon exponentielle avec la tension extérieure

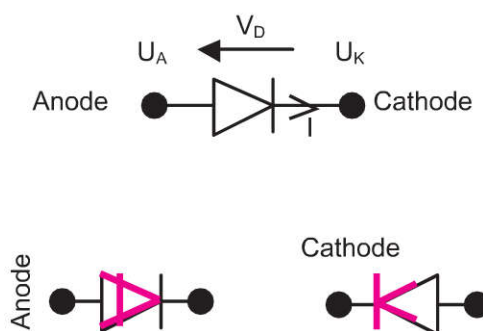
c) En pratique

La jonction PN simple permet la création de composants électroniques : les diodes à jonction. Son symbole, asymétrique, indique que ces composants sont polarisés : leur comportement n'est pas le même selon le sens dans lequel ils sont branchés.

Le sens de la diode peut se déduire aisément à partir de son symbole : l'anode est du côté où le symbole forme un A couché, et la cathode (Kathode en allemand) est du côté où le symbole forme un K.

Lorsque nous parlerons de la tension V_D appliquée à la diode, celle-ci sera égale à la différence de potentiel $U_A U_K$.

En polarisation directe, $U_A > U_K$, $V_D > 0$, la diode est passante.



$$V_D = \frac{k \cdot T}{q} \cdot \ln \left(\frac{N_A \cdot N_D}{n_i^2} \right)$$

T : température en Kelvin

q : charge de l'électron

k : constante de Boltzmann ($1,38 \times 10^{-23} \text{ J} \cdot \text{K}^{-1}$)

N_A : densité des atomes accepteurs, c'est-à-dire du dopage P.

N_D : densité des atomes donneurs, c'est-à-dire du dopage N.

n_i : densité de porteurs libres dans le semi-conducteur intrinsèque (non dopé).

À la température de 300 K proche de l'ambiante, $\frac{k \cdot T}{q}$ vaut environ 26 mV.

$$I_D = I_S \cdot \left(e^{\frac{q \cdot V_D}{k \cdot T}} - 1 \right)$$

I_S est le courant de saturation dû aux porteurs minoritaires. Ce courant dépend de la température. En polarisation inverse, $U_A < U_K$, $V_D < 0$, la diode est bloquée.

d) Caractéristiques électriques

	1	Zone d'avalanche : en-deçà de V_Z , des électrons fortement accélérés franchissent la barrière de potentiel en inverse et entraînent d'autres électrons arrachés au réseau cristallin.
	2	Blocage : lorsque la tension d'entrée est inférieure à V_Y , seul le courant inverse très faible I_S circule.
	3	Conduction : lorsque la tension d'entrée est supérieure à V_Y , la diode est passante, on peut appliquer la formule $I_D = f(V_D)$.

e) Approximations de la diode : modèles équivalents

Modèles	Caractéristiques électriques	Schémas équivalents
Diode idéale $V_Y = 0 \text{ V}$		$V_D < 0$ $V_D > 0$
1^{re} approximation $V_Y > 0 \text{ V}$		$V_D < V_Y$ $V_D > V_Y$
Approximation linéaire		$V_D < V_Y$ $V_D > V_Y$
Approximation exponentielle		$V_D < V_Y$ Aucun $V_D > V_Y$

31 Différents types de diodes

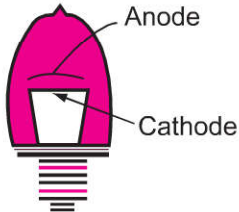
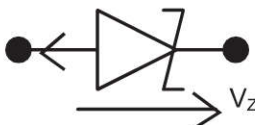
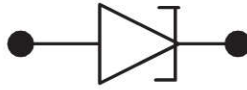
Mots clés

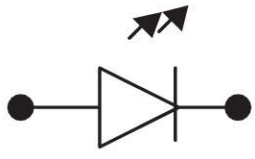
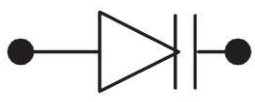
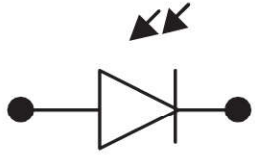
LED, Zener, Schottky, PIN, photodiodes, redressement simple alternance.

1. EN QUELQUES MOTS :

Nous avons vu dans la **fiche 30** que la diode est un composant électronique (ou un dipôle électrique monodirectionnel) non-linéaire fabriqué à partir de semi-conducteurs de types P et N. Le courant circule de l'anode (potentiel U_A) vers la cathode (potentiel U_K) : en première approximation, si $U_A > U_K$, la diode se comporte comme un interrupteur fermé ou un fil, et pour $U_A < U_K$, la diode est équivalente comme un interrupteur ouvert. Selon ce principe de base, différents types de diodes existent et trouvent leur application pour différents montages.

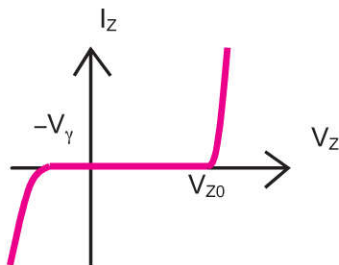
2. CE QU'IL FAUT RETENIR

Diode	Fonctionnement	Symbole
de Fleming	Première diode construite en 1903, elle comporte une cathode chargée des électrons lorsqu'elle est chaude, une anode (cylindre de tôle mince qui entoure la cathode), un filament alimenté par un courant électrique en basse tension.	
Zener	La diode Zener accepte une tension inverse (tension Zener) ou tension d'avalanche de valeur déterminée de 1,2 V à plusieurs centaines de volts. Cette valeur est réglable sur certains modèles.	
Schottky	La diode Schottky est réalisée à partir d'une jonction métal-semi-conducteur qui a un seuil de tension directe très bas et un temps de commutation très court.	
à effet tunnel	Une diode à effet tunnel est rapide et stable. On l'utilise souvent dans les circuits où un temps de commutation très court est indispensable. Elle est utilisée à basse puissance.	
laser	Une diode laser est un composant optoélectronique. Exemples d'application : système d'information optique, barrage photoélectrique avec grande portée, appareil de mesure analytique.	
PIN	Une diode PIN (<i>Positive Intrinsic Negative</i>) est constituée de deux zones dopées P et N et d'une zone intrinsèque I au centre. Elle présente une commutation rapide.	

Diode	Fonctionnement	Symbole
électroluminescente	Une diode électroluminescente est une diode lumineuse : on l'appelle aussi DEL ou LED. Le courant électrique passe uniquement dans un seul sens et produit un rayonnement monochromatique incohérent.	
varicap	Diode utilisée comme condensateur variable dans les circuits des récepteurs radios et des téléviseurs.	
à vapeur de mercure (Tube électronique)	La diode à mercure est constituée d'un tube électronique comme les diodes de Fleming. On l'utilise dans les thermomètres électroniques.	
Gunn	Une diode Gunn est constituée de trois régions de type N. On l'utilise en électronique d'extrêmement hautes fréquences : elle permet la construction d'oscillateurs à partir de 8 GHz.	
Photodiode	Les photodiodes ont la capacité de détecter un rayonnement optique et de le transformer en signal électrique. Lorsque la diode reçoit de la lumière, elle reçoit de l'énergie. Par conséquent, le nombre de porteurs minoritaires augmente et I_s augmente. Elles sont polarisées en inverse : courant inverse de saturation du μA à quelques mA dû aux porteurs minoritaires.	
Transil	Dans cette diode, on applique le même effet d'avalanche que la diode Zener, mais dans un but exclusif de protection des circuits. Les diodes Transil présentent les avantages suivants : plus faible résistance dynamique, meilleure précision et plus grande dynamique dans le choix de la tension d'avalanche, meilleure fiabilité.	

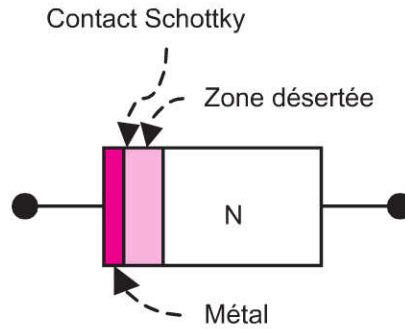
a) Quelques précisions

Une **diode Zener** est une diode que l'on fait fonctionner dans la zone d'avalanche.

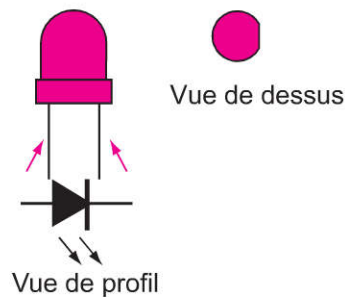


Par convention, on pose
$$\begin{cases} I_Z = -I_D \\ V_Z = -V_D \end{cases}$$

- La **diode Schottky** est un composant unipolaire. Les seuls porteurs libres sont les électrons.



- Les **diodes électro-luminescentes** ou DEL (en anglais : *Light Emission Diode* – DEL) les plus couramment employées aussi bien au niveau grand public que professionnel sont du type 1N400X. Elles supportent 1 A. Elles peuvent être au format CMS ou traversant comme l'indique la figure.

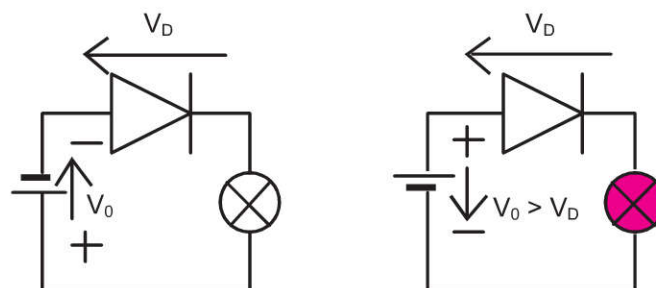


3. EN PRATIQUE

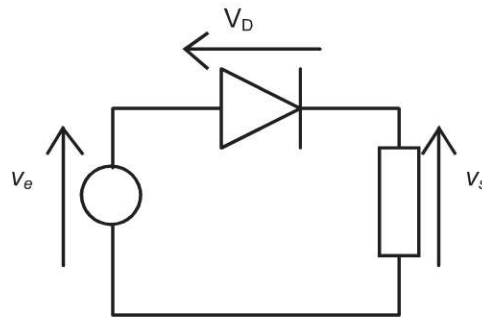
Les applications des diodes sont nombreuses : redressement de tension, doubleurs de tension, détection de signaux RF dans un récepteur radio, protection contre les surtensions, amélioration de la vitesse de commutation de transistors de puissance dans une alimentation à découpage et dans la fabrication de certains interrupteurs.

Le **redressement simple alternance** consiste à éliminer les valeurs négatives d'une source de tension. Dans le cas simple constitué d'une pile, d'une diode et d'une lampe, la lampe s'éclaire dans la situation du schéma de droite mais pas dans celle de gauche :

- à gauche, la diode est polarisée en inverse, il n'y a aucun courant parcourant le circuit ;
- à droite, la diode est polarisée en direct. Si la tension d'alimentation continue V_D est supérieure à la tension de seuil de la diode, un courant circule et la lampe éclaire.

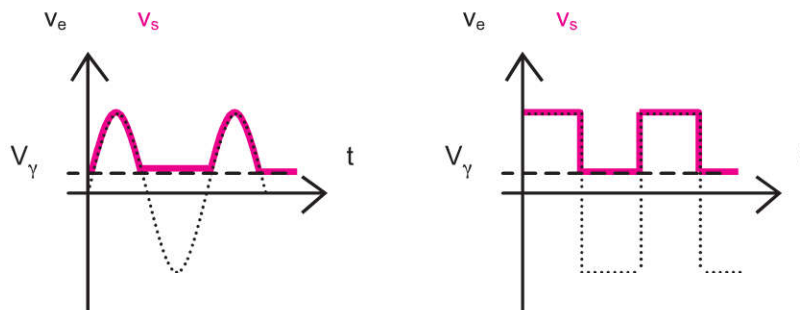


Si maintenant, nous remplaçons la source de tension continue par un générateur de tension alternative (de forme *a priori* quelconque), seules les tensions supérieures à la tension de seuil de la diode seront transmises à la résistance de charge R.

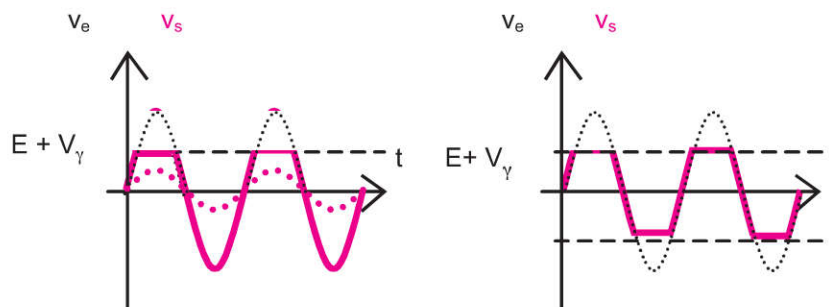
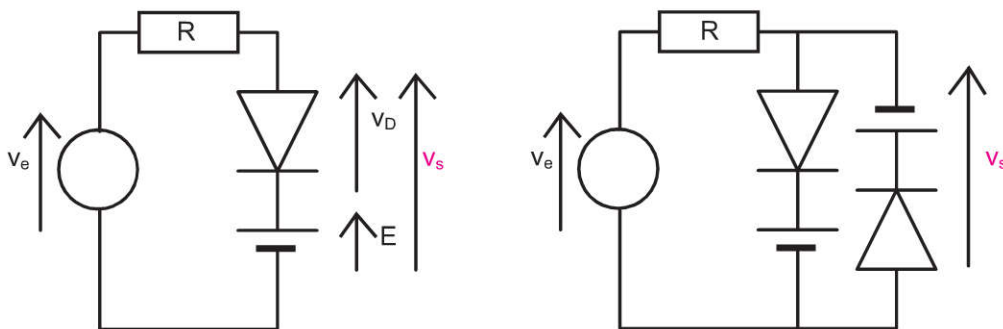


$$\left. \begin{array}{l} v_s = v_e - V_D \\ v_s = R \cdot i_D \end{array} \right\} i_D = \frac{v_e - v_d}{R}$$

La diode conduit si la tension d'entrée v_e est supérieure à la tension de seuil.



Le cas de l'**écrêteur-limiteur** est assez similaire au redresseur de tension : on positionne une source de tension continue en série avec une diode afin de régler le niveau de tension v_s maximum disponible. Ce dispositif peut être répliqué et monté tête-bêche afin d'écarter la tension en négatif et en positif (figure de droite).



Le **redressement double alternance** est abordé dans la [fiche 72](#).

32 Transistor bipolaire

Mots clés

Jonction, effet transistor, base, collecteur, émetteur, commutation, linéaire.

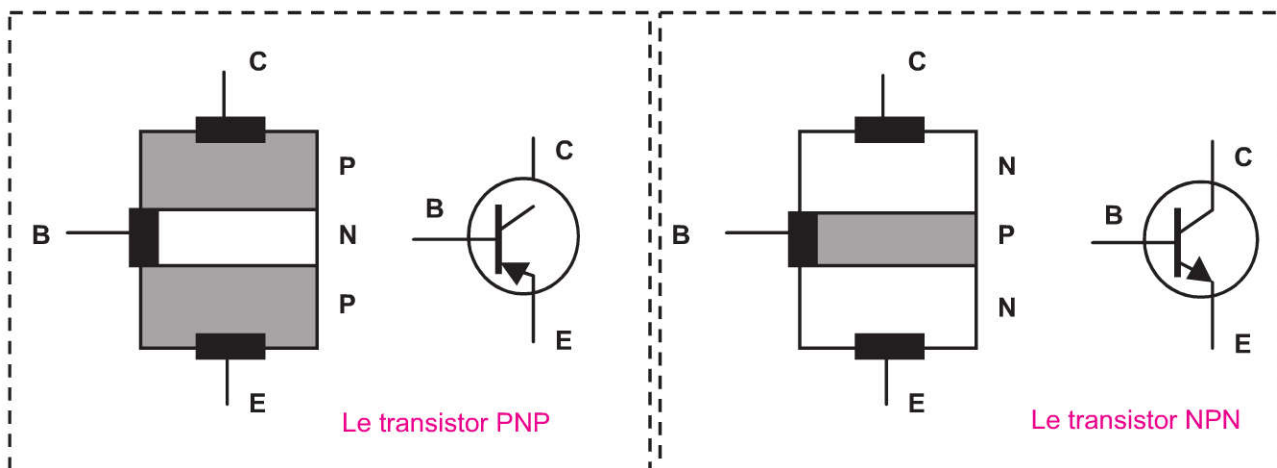
1. EN QUELQUES MOTS

Deux types de transistors bipolaires existent : le transistor de type PNP et NPN. Il est constitué de trois broches : la base (B), le collecteur (C) et l'émetteur (E).

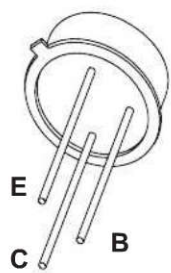
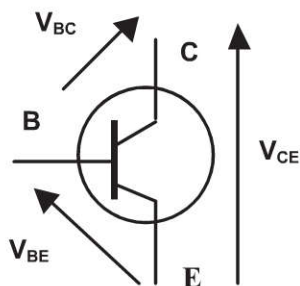


2. CE QU'IL FAUT RETENIR

La représentation physique et le symbole de chacun d'entre eux peut être décrit par les figures ci-dessous :



On notera par convention les tensions comme des différences de potentiels pour que $V_{BE} = V_B - V_E$.



Package TO18

Ainsi nous définissons les tensions :

entre base-émetteur: V_{BE}

entre collecteur-base: V_{CB}

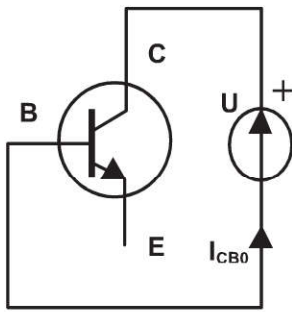
entre collecteur-émetteur : V_{CE}

De même pour les courants notés I_B , I_C et I_E .

► Fonctionnement non polarisé

La description statique de ce composant permet de mettre en évidence qu'un transistor est un modèle à deux diodes mises en série (mais avec des caractéristiques qui ne sont pas symétriques) avec anodes ou cathodes communes par la base. L'explication est donnée sur une base de transistor de NPN.

► Fonctionnement à émetteur non connecté



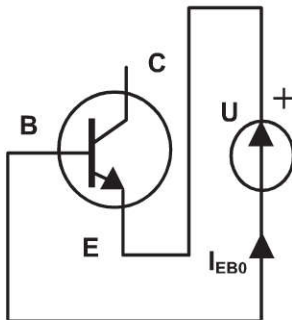
Contexte :

- L'émetteur n'est pas connecté.
- Une tension U positive est appliquée en direct entre la base et le collecteur.

Un courant circule entre ces deux broches, il est noté conventionnellement I_{CB0} .

Exemple : valeur pour un transistor de type NPN 2N2222A, $I_{CB0} = 10\text{ nA}$ avec $U = V_{CB} = 60\text{ V}$ avec forcément $I_E = 0\text{ mA}$.

► Fonctionnement à collecteur non connecté



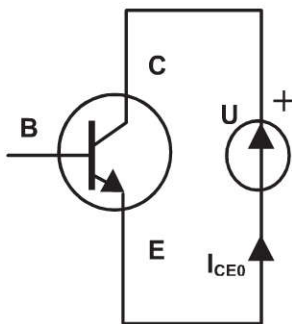
Contexte :

- Le collecteur n'est pas connecté.
- Une tension U positive est appliquée en direct entre la base et l'émetteur.

La première lettre du courant correspond au courant entrant dans la broche

Exemple : valeur pour un transistor de type NPN 2N2222A, $I_{EB0} = 10\text{ nA}$ avec $U = V_{EB} = 3\text{ V}$.

► Fonctionnement à base non connectée



Contexte :

- La base n'est pas connectée.
- Une tension U positive est appliquée en direct entre l'émetteur et le collecteur.

L'indice 0 du courant indique que la lettre n'apparaissant pas (ici B) est non connectée.

a) Effet transistor

Lors d'une polarisation directe d'un transistor PNP, par exemple, pour qu'un nombre maximal d'électrons passe de l'émetteur au collecteur, il faut que l'une des deux conditions soit remplie :

- la base doit être fine et la surface de la jonction collecteur-base importante ;
- la base est faiblement dopée et, à l'inverse, l'émetteur fortement dopé.

La structure ainsi réalisée, le pourcentage de recombinaisons est très faible. La grande majorité des électrons émis depuis l'émetteur est reçue par le collecteur. Ce phénomène est appelé **effet transistor**.

Quels que soient les types de montages étudiés, les six variables du transistor sont liées par la loi des nœuds et des mailles :

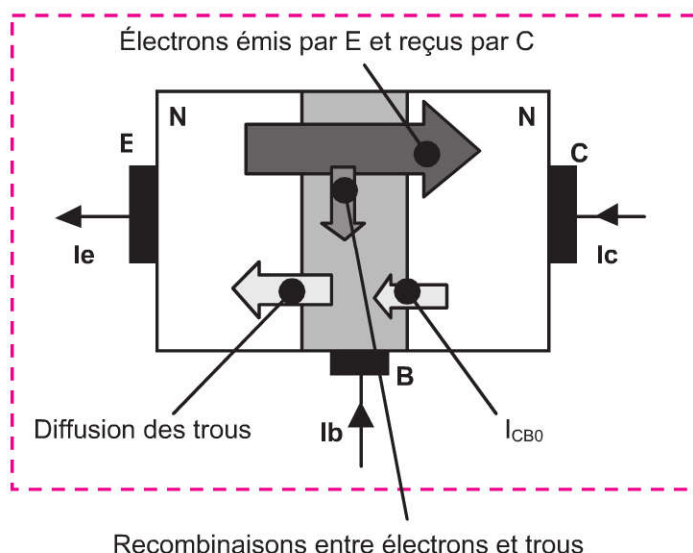
$$\begin{cases} I_b + I_e + I_c = 0 \\ V_{ce} + V_{eb} + V_{bc} = 0 \\ I_c = \beta \times I_b \\ I_e = I_c + I_b = (\beta + 1) \times I_b \approx \beta \times I_b \end{cases}$$

où β représente le gain en courant du transistor. Pour un 2N2222, il apparaît sous le nom de h_{FE} sur la fiche constructeur, dépendant des conditions sur les valeurs de I_C et V_{CE} . Ce paramètre fluctue beaucoup en fonction de la construction et il vaut au moins 30 pour ce modèle. Pour d'autres, il peut atteindre 300 selon le type de transistor (souvent pour les transistors basse puissance < 200 mW).

3. EN PRATIQUE

a) Régimes de fonctionnement

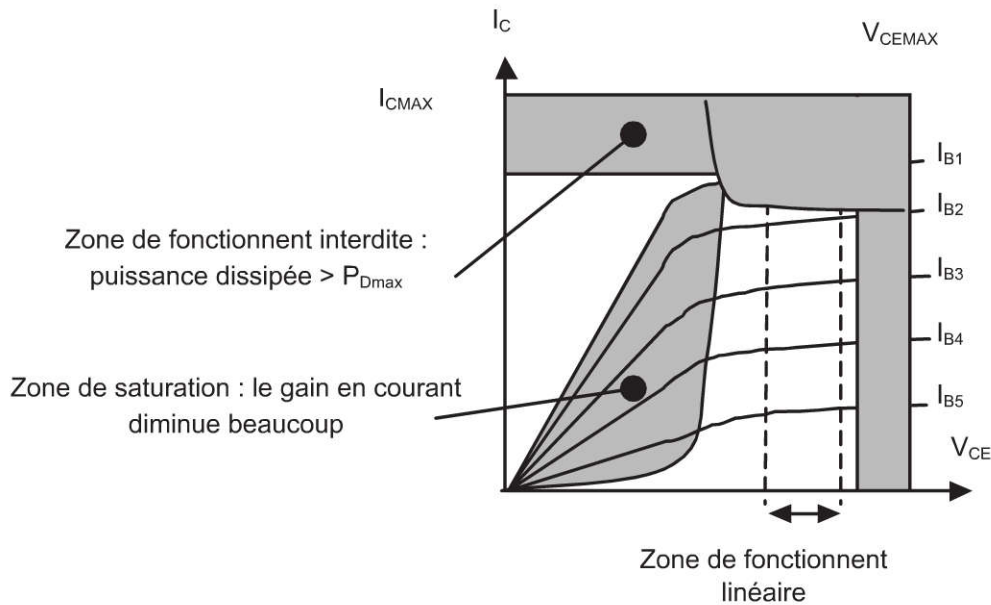
Le **fonctionnement en commutation**, ou mode interrupteur, est à la base de la construction des **circuits numériques** : le transistor se comporte comme un interrupteur dont les deux états peuvent représenter un 1 ou un 0 binaire.



Type et état	Conditions
	NPN Bloqué ou ouvert Si $I_B = 0$ alors : - $V_{CE} > 0$ - V_{CE} doit être $< V_{CEmax}$ - $I_C = 0$ - $I_C > 0$ par rapport à la convention
	NPN Saturé ou fermé Si $I_B = I_{BSAT} > 0$ alors : - $V_{CE} = V_{CEsat} > 0$ - I_C doit être $< I_{Cmax}$ - $I_C = \beta \times I_{BSAT} > 0$ - $I_C > 0$ par rapport à la convention
	PNP Bloqué ou ouvert Si $I_B = 0$ alors : - $V_{CE} < 0$ par rapport à la convention - V_{CE} doit être $< V_{CEmax}$ - $I_C = 0$ - $I_C < 0$ par rapport à la convention
	PNP Saturé ou fermé Si $I_B = I_{BSAT} > 0$ alors : - $V_{CE} = V_{CEsat} < 0$ par rapport à la convention - I_C doit être $< I_{Cmax}$ - $I_C = \beta \times I_{BSAT} < 0$ - $I_C < 0$ par rapport à la convention

Pour être certain de saturer le transistor, on calcule $I_B = k \times I_{CSAT} / \beta$ avec $k > 3$ et $I_B < I_{Bmax}$ du constructeur.

On utilise la caractéristique de sortie du transistor dans sa zone linéaire $V_{CE} = f(I_C)$. Pour un montage dont l'émetteur est au commun pour un type NPN nous obtenons le type de caractéristique suivant :



Cette caractéristique est paramétrée par I_B ici $I_{B1} > I_{B2} > \dots > I_{B5}$. Pour avoir un ordre d'idée, la zone de fonctionnement linéaire est de l'ordre de 1 à quelques Volts pour V_{CE} . En dehors de cette zone, il s'agit d'un fonctionnement de type saturé qui nous amène au mode interrupteur décrit précédemment.

La zone de fonctionnement linéaire est utilisée pour amplifier des signaux émis depuis un courant de base (cas des montages à émetteur commun). On peut séparer plusieurs zones de fonctionnement mais elles ne peuvent se chevaucher. Le résultat d'un chevauchement se traduirait par une non-linéarité de la caractéristique de sortie.

Les signaux sont de faible amplitude comparés aux signaux d'alimentation du montage (environ $< 5\%$).

b) En résumé

Un transistor bipolaire peut être utilisé en interrupteur (tout ou rien). On dit alors que :

- le transistor est **saturé** lorsqu'il est passant ou dit fermé ;
- le transistor est **bloqué** lorsqu'il est ouvert.

Le transistor bipolaire peut être utilisé en amplification. Dans ce cas, c'est le courant I_B (cas d'un montage en émetteur commun (fiche 37)) qui permettra de faire varier le courant I_C autour d'un point de repos. Le point de repos dépendra du montage de polarisation. (fiche 33)

La valeur de β peut varier dans un rapport 1 de 3 pour des transistors de la même famille. Par ailleurs, β est fonction de la température. Ce qui implique que chaque montage est différent même si la référence du transistor est identique.

Dans tous les cas, pour un transistor bipolaire, c'est **un courant qui permet de régler un autre courant**.

33 Réseau de caractéristiques

Mots clés

Caractéristiques d'entrée, de sortie, de transfert ; point de fonctionnement ; droite de charge et d'attaque.

1. EN QUELQUES MOTS

Nous avons vu dans la [fiche 32](#) que les transistors bipolaires présentent différents modes de fonctionnement. Outre les équations, une analyse des réseaux de caractéristiques permet de comprendre comment polariser et utiliser ces composants.

2. CE QU'IL FAUT RETENIR

a) Réseau de caractéristiques : cas particulier d'un NPN avec Émetteur au commun

► Montage associé mettant en évidence le réseau de caractéristiques

Pour expliquer la nature de ce réseau, commençons par étudier le montage le plus souvent utilisé, à savoir l'utilisation d'un transistor NPN avec l'émetteur connecté à la masse (émetteur commun).

Le montage permettant la commande du transistor (montage d'entrée) est constitué d'une source de tension variable E_E et d'une résistance R_B qui permet de limiter le courant I_B dans la base.

Le montage de sortie est similaire au montage d'entrée, avec une source de tension variable E_S et une résistance R_C .

Pour ce montage, les caractéristiques :

- d'entrée sont I_B et $V_{BE} \rightarrow V_{BE} = f(I_B)$ à $V_{CE} = \text{cst}$;
- de sortie sont I_C et $V_{CE} \rightarrow I_C = f(V_{CE})$ à $I_B = \text{cst}$.

► Caractéristique de sortie : $I_C = f(V_{CE})$ à $I_B = \text{cst}$

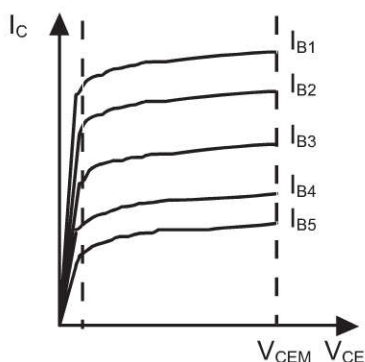
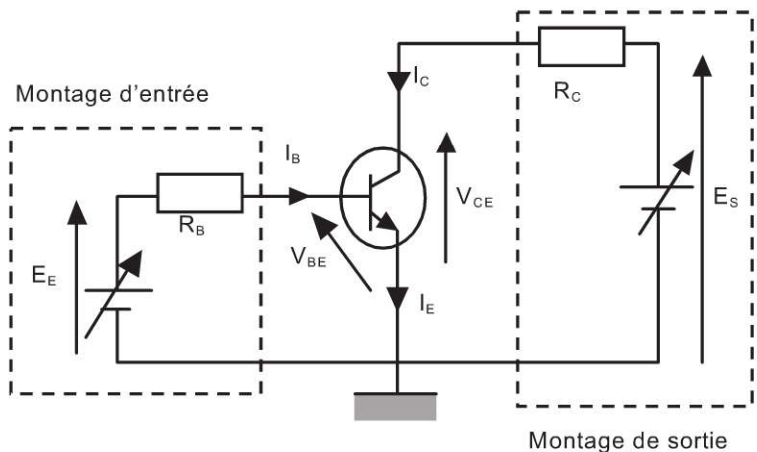
Conditions

Pour une valeur de E_E qui règle une valeur I_B , on fait varier la valeur de E_S de 0 à $E_{S\text{MAX}}$ et l'on mesure V_{CE} et I_C . Puis on réitère cette opération pour plusieurs valeurs de E_E donc de I_B .

On obtient alors le réseau de courbes ci-contre.

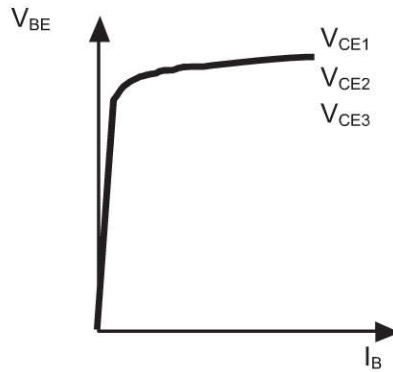
$I_{B1} > I_{B2} > \dots > I_{B5}$, il est à noter que I_B ne doit pas dépasser le I_{BM} du constructeur et il en va de même pour V_{CEM} .

La zone où V_{CE} est proportionnelle à I_C est grande comparée à la zone saturée.



► **Caractéristique d'entrée :** $V_{BE} = f(I_B)$ à $V_{CE} = \text{cst}$

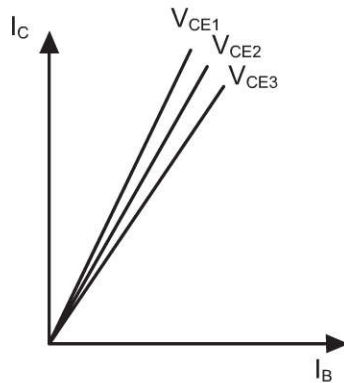
Conditions



Pour une valeur de E_S qui règle une valeur V_{CE} , on fait varier la valeur de E_E de 0 à E_{EMAX} et l'on mesure V_{BE} et I_B . Puis on réitère cette opération pour plusieurs valeurs de E_S donc de V_{CE} . On obtient alors le réseau de courbes ci contre où toutes les courbes se superposent, ce qui est homogène avec le modèle du transistor. Sachant que V_{BE} n'est autre que la tension présente aux bornes d'une jonction P-N (base-émetteur). L'ordre de grandeur de cette tension est de 0,2 V pour une jonction au germanium à 1 V pour une jonction silicium.

► **Caractéristique transfert en courant (entrée – sortie) de:** $I_C = f(I_B)$ à $V_{CE} = \text{cst}$

La relation entre l'entrée et la sortie définit la **caractéristique de transfert en courant**. Cette dernière est linéaire. C'est la pente qui définit le **gain en courant** du transistor, appelée h_{FE} ou β dans les notices constructeurs. La valeur du gain est légèrement dépendante de la valeur de V_{CE} ($V_{CE1} > V_{CE2} > V_{CE3}$). Les constructeurs garantissent un h_{FE} min pour des conditions spécifiées.



$$h_{FE} = I_C / I_B$$

Remarque

Il est possible de définir aussi une caractéristique entrée-sortie $V_{BE} = f(V_{CE})$ à $I_B = \text{cst}$ mais cette dernière n'est quasiment jamais exploitée.

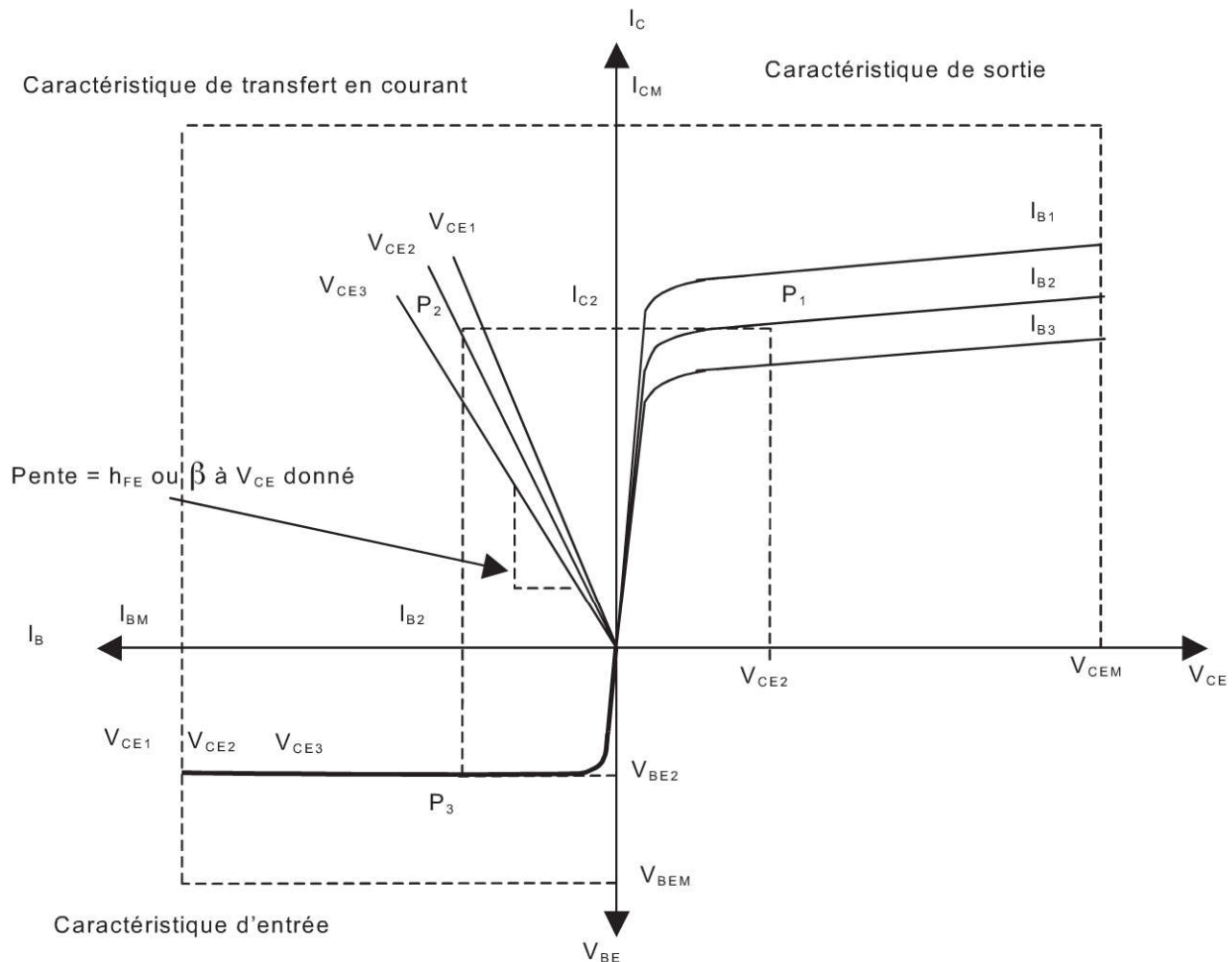
► **Synthèse des caractéristiques**

Les trois caractéristiques précédentes peuvent être regroupées dans un repère de quatre quadrants. Ces derniers font correspondre chacune des grandeurs avec les autres.

Ainsi, en partant d'une valeur V_{CE2} , nous arrivons en P_1 nous indiquant, à valeur de I_B donnée (I_{B2} dans notre cas), la valeur de I_{C2} . En prolongeant, nous obtenons le point P_2 à valeur de V_{CE}

donnée (ici V_{CE2}). Le dernier point P_3 , par la valeur de son abscisse, nous confirme la valeur de I_B (ici I_{B2}), et par la valeur de son ordonnée, la valeur de V_{BE} (ici V_{BE2}).

On peut y faire apparaître les valeurs maximales caractéristiques du transistor utilisé (V_{CEM} , I_{CM} ...). Il est important aussi de préciser que les échelles de chaque axe ne sont pas du même ordre de grandeur (exemple : $V_{BE} < 1$ volt alors que V_{CE} peut être de plusieurs dizaines de volt).



b) Droite d'attaque et de charge : cas particulier d'un NPN avec émetteur au commun

Toujours en partant du même montage, les valeurs de R_C et R_B à E_S et E_E données fixent un point de fonctionnement sur le réseau de caractéristiques.

Plaçons-nous du côté montage de sortie : nous devons exprimer V_{CE} en fonction de E_S , I_C et R_C . D'où $V_{CE} = E_S - R_C \times I_C$.

Du côté montage d'entrée, il s'agit de V_{BE} fonction de E_E , I_B et R_B d'où $V_{BE} = E_E - R_B \times I_B$.

- On appelle **droite d'attaque** la caractéristique $I_B = f(V_{BE})$ en fonction des valeurs de composants et source côté « montage d'entrée ». Dans notre cas, $I_B = (E_E - V_{BE})/R_B$.
- De même, on désigne **droite de charge** la caractéristique $I_C = f(V_{CE})$ à I_B donnée, par la caractéristique d'attaque en fonction des valeurs de composants et source coté « montage de sortie », dans notre cas $I_C = (E_S - V_{CE})/R_C$.

Exemple

$E_E = E_S = 15\text{ V}$, $R_C = 500\ \Omega$, $R_B = 100\text{ k}\Omega$, $h_{FE} = 100$ que l'on supposera constant quel que soit V_{CE} , $V_{CEM} = 30\text{ V}$, $I_{BM} = 5\text{ mA}$, $I_{CM} = 0,5\text{ A}$.

Pour la droite d'attaque :

Si $I_B = 0$ alors $V_{BE} = E_E = 15\text{ V}$,

et si $V_{BE} = 0$ alors $I_B = E_E / R_B = 15 / 100 \times 10^3 = 0,15\text{ mA}$.

Cela donne les deux points permettant le tracé de la droite. L'intersection de la droite d'attaque avec la caractéristique $I_B = f(V_{BE})$ donne le point P_3 .

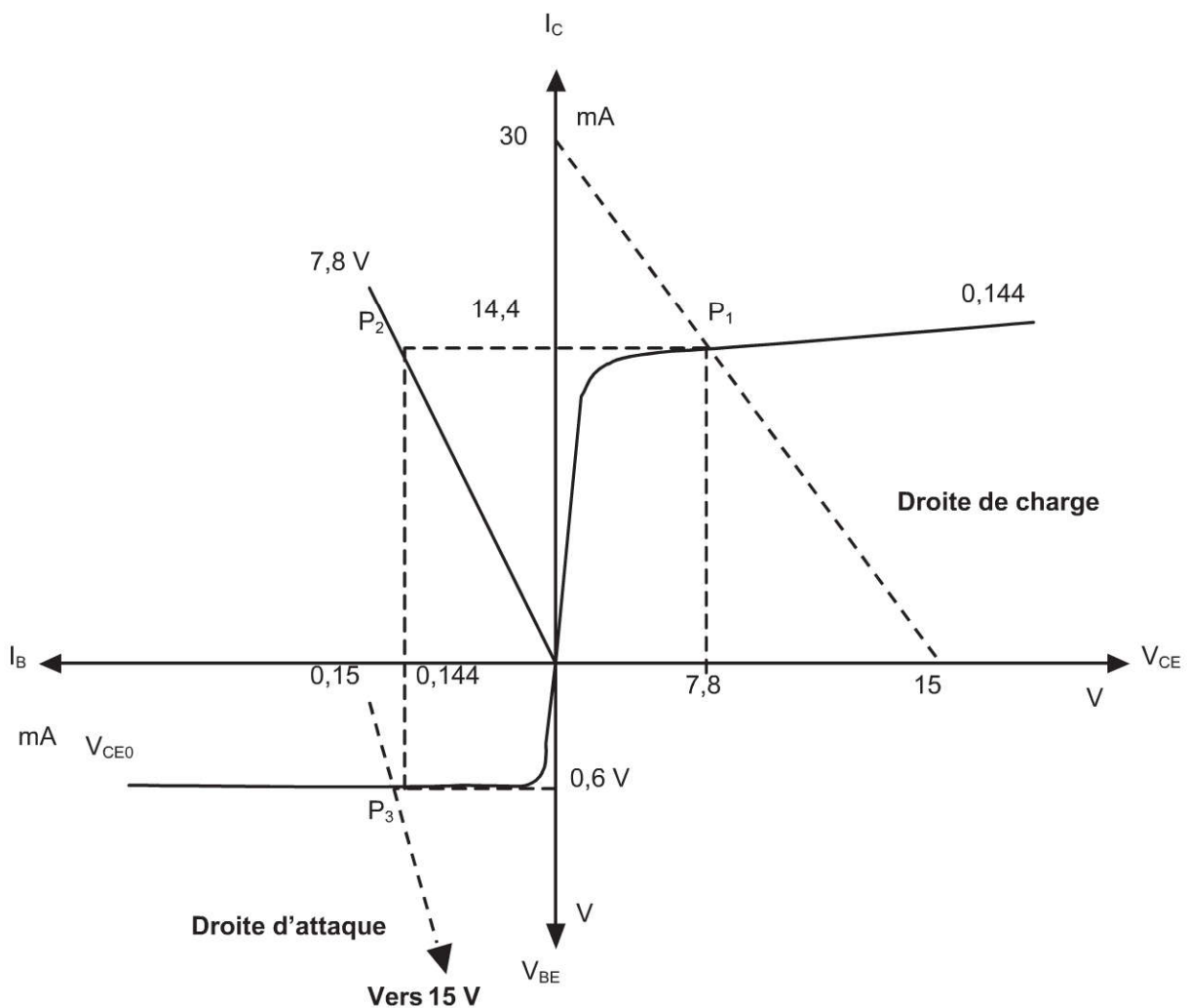
Le point P_3 peut être aussi calculé : $I_B = (15 - 0,6) / 100 \times 10^3 = 0,144\text{ mA}$ donne le point P_3 .

Pour la droite de charge :

Si $I_C = 0$ alors $V_{CE} = E_S = 15\text{ V}$, et si $V_{CE} = 0$ alors $I_C = E_S / R_C = 15 / 500 = 30\text{ mA}$.

La détermination des coordonnées du point P_1 sont déduites de I_B calculé du point P_3 ($0,144\text{ mA}$), avec un $\beta = 100$ on obtient $I_C = 100 \times 0,144 = 14,4\text{ mA}$.

La tension $V_{CE} = E_S - R_S \times I_C = 15 - 500 \times 14,4 \times 10^{-3} = 7,8\text{ V}$

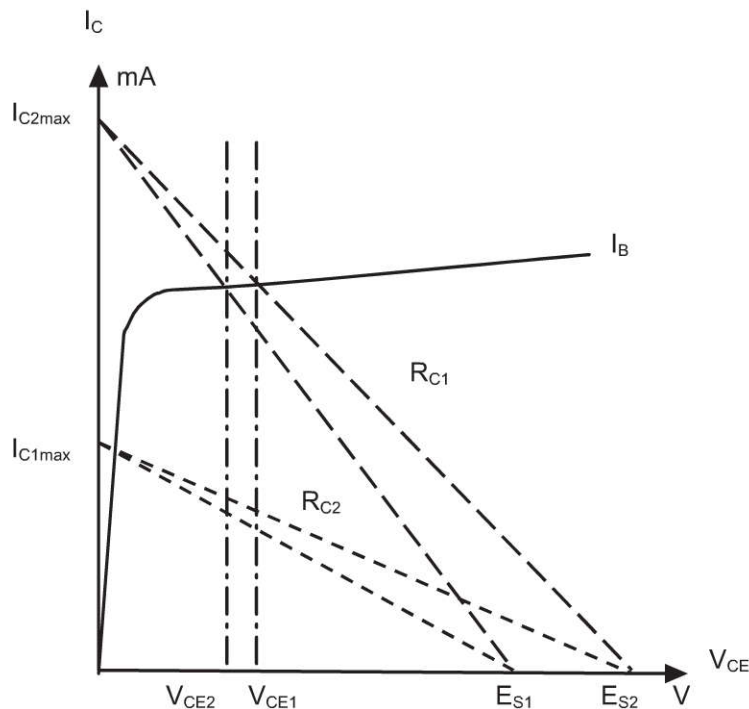


L'intersection de la droite de charge avec la caractéristique $I_C = f(V_{CE})$ donne le point P_1 .

3. EN PRATIQUE

a) Modification du point de repos

► Pente de la caractéristique de charge à I_B donné

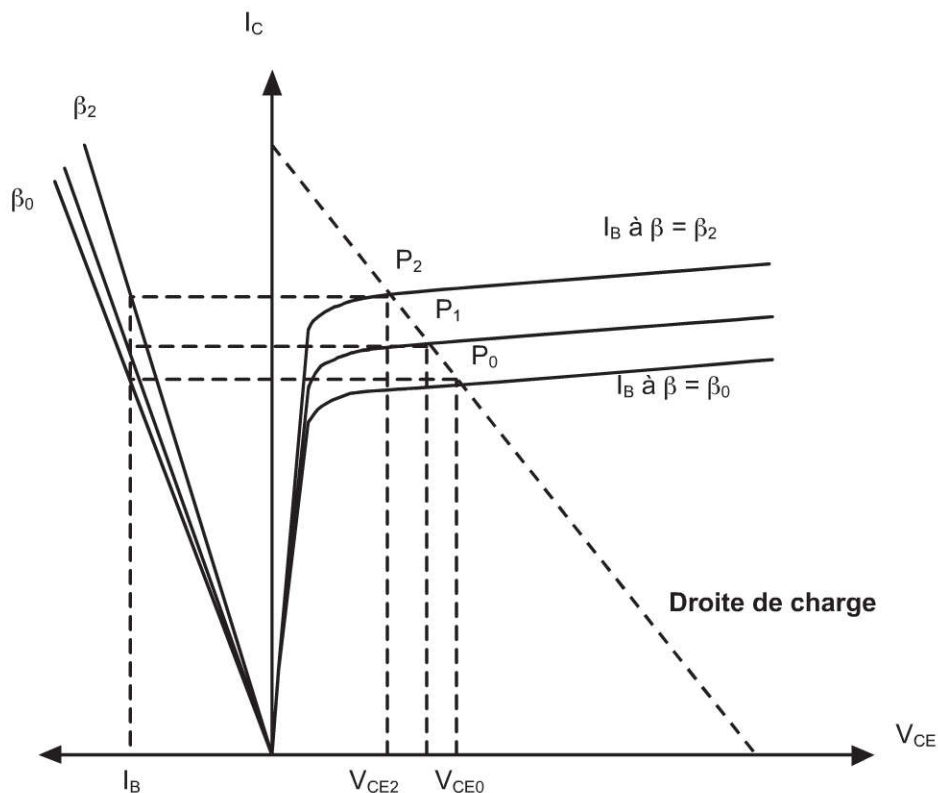


Elle dépend forcément du montage de sortie et donc de R_C et E_S . En notant R_{C1} et R_{C2} , deux résistances du montage de sortie différentes avec $R_{C1} < R_{C2}$, ainsi que E_{S1} et E_{S2} , deux valeurs de tensions E_S avec $E_{S1} < E_{S2}$, on peut observer que les deux paramètres conjugués permettent d'obtenir quatre pentes de droites de charge différentes.

À E_S donné, si R_C augmente, I_{Cmax} diminue, la pente diminue. De même, à R_C fixé, si on augmente la tension E_S , I_{Cmax} augmente, la pente augmente aussi. Pour la même valeur de I_B , on obtient alors des valeurs de V_{CE}

différentes (Exemple : V_{CE1} et V_{CE2}).

► Changement du point de repos à E_S et R_C donnés



La valeur de h_{FE} (ou β) est quasiment indépendante de V_{CE} ainsi, seul le montage d'entrée fixe la valeur de I_B donc de I_C .

Si le gain vient à être différent pour un transistor d'une même série, alors le point de fonctionnement sera déplacé de P_1 vers P_0 si le gain diminue (β_0) et de P_1 à P_2 dans le cas inverse (β_2).

► **Le point de fonctionnement impossible**

Des incohérences peuvent apparaître si le choix de E_S , E_E et du transistor n'est pas harmonieux. Par exemple, fixons $E_E = 10\text{ V}$, $E_S = 20\text{ V}$, $R_B = 100\text{ k}$, $R_C = 10\text{ k}$, β mini du transistor indiqué sur la notice constructeur vaut 50.

$$I_B = (E_E - V_{BE}) / R_B = (10 - 0,6) / 100 \cdot 10^3 = 0,1\text{ mA}$$

$$I_C = \beta \times I_B = 50 \times 0,1 = 5\text{ mA}.$$

Si le transistor est complètement saturé $V_{CE} \cong 0$ et $I_{C_{\max}} = E_S / R_C = 20 / 10 \cdot 10^3 = 2\text{ mA}$, non cohérent avec la valeur de I_C calculée... I_C ne peut être supérieure à $I_{C_{\max}}$.

Autre façon de voir, $V_{CE} = E_S - R_C \times I_C = 20 - 10 \cdot 10^3 \times 5 \cdot 10^{-3} = -30\text{ V}$, ce qui est impossible.

b) En résumé

- Le transistor est un quadripôle commandé en courant par un montage d'entrée pouvant être décrit par la droite d'attaque. Le montage de sortie peut être caractérisé par la droite de charge.
- Le montage d'entrée a pour paramètres E_E et R_B qui fixent I_B et le montage de sortie a pour paramètres E_S et R_C .
- I_B fixe la valeur de I_C .
- La droite de charge est utile pour visualiser tous les points de fonctionnement réalisables par le montage à valeur de I_B fixé.
- Les choix des valeurs E_E , E_S , R_B et R_C ne peuvent être arbitraires mais doivent être compatibles avec les caractéristiques du transistor.
- Les valeurs de I_C et V_{CE} étant en partie dépendantes de β , on dit que ce montage est β dépendant.
- Relations :
 - pour le montage d'entrée pour un NPN à émetteur au commun :

$$I_B = (E_E - V_{BE}) / R_B$$
 - Pour le montage de sortie pour un NPN à émetteur au commun :

$$V_{CE} = E_S - R_C \times I_C$$
 - Entre le montage d'entrée et de sortie : $I_C = \beta \times I_B$

34 Polarisation d'un transistor

Mots clés

Point de fonctionnement, dépendant β , indépendant β .

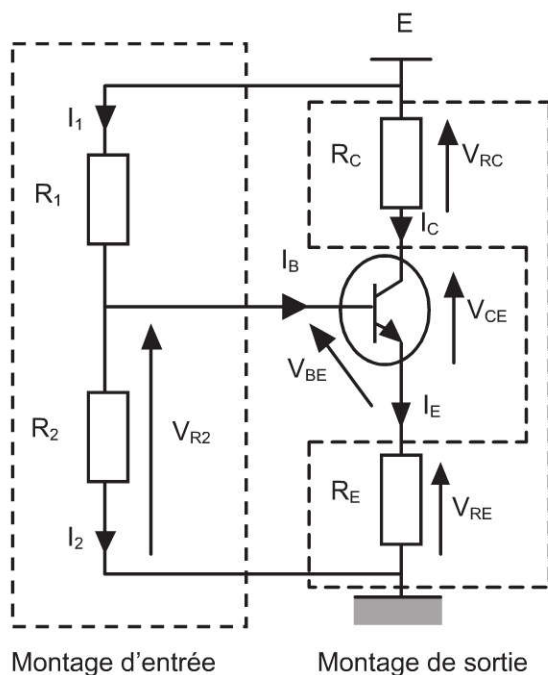
1. EN QUELQUES MOTS

Afin de pouvoir utiliser un transistor, quel qu'il soit, il est indispensable de le polariser, ce qui se traduit, en d'autres termes, par imposer un point de repos. Cela consiste à fixer les valeurs continues de tensions et de courants autour desquelles les grandeurs alternatives peuvent évoluer.

2. CE QU'IL FAUT RETENIR

a) Polarisation en pont d'un transistor NPN

Il est nécessaire de préciser que les valeurs de R_1 et R_2 sont définies de façon à ce que I_1 , I_2 et donc I_B soit négligeable devant I_2 . Négligeable signifie que $I_2 / I_B > 100$. Cela permet d'utiliser la relation du diviseur de tension.



$$V_{R2} = \frac{R_2}{R_2 + R_1} \times E$$

On reconnaît une polarisation par la base avec une tension de commande du montage d'entrée $E_E = V_{R2}$. On peut en déduire que :

$$\begin{cases} V_{RE} = V_{R2} - V_{BE} \\ V_{RE} = R_E \times I_E \\ I_E = I_C + I_B = \beta \times I_B + I_B = (\beta + 1) \times I_B \approx I_C \\ V_{RC} = E - R_C \times I_C \\ V_{CE} = E - V_{RC} - V_{RE} \end{cases}$$

► Méthode applicative pour obtenir les valeurs des résistances

On impose $E = 10 \text{ V}$, $\beta_{\text{mini}} = 100$ de la notice constructeur, $I_C = 10 \text{ mA}$, $V_{BE} = 0,7 \text{ V}$. On s'impose le V_{CE} du point de repos à une valeur de 50 % de E et on répartit $V_{RE} = 10 \%$ de E et le reste à V_{RC} soit 40 % de E .

$$V_{RE} = 0,1 \times 10 = 1 \text{ V}, V_{RC} = 0,4 \times 10 = 4 \text{ V} \text{ et } V_{CE} = 5 \text{ V}$$

$$I_C \approx I_E = 10 \text{ mA d'où } R_E = V_{RE} / I_E = 1 / 0,01 = 100 \Omega \text{ de même } R_C = 4 / 0,01 = 400 \Omega$$

$I_B = I_C / \beta = 10 \text{ mA} / 100 = 0,1 \text{ mA}$ et donc par hypothèse I_2 doit être 100 fois plus grand $I_2 = 10 \text{ mA}$. Ce qui vérifie $I_1 \cong I_2$.

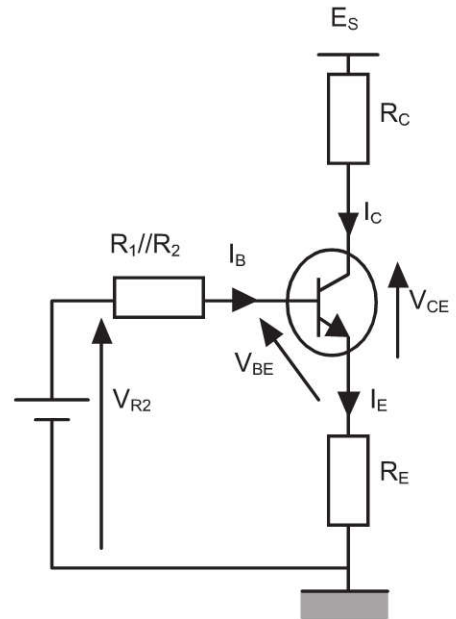
$$V_{R2} = V_{RE} + V_{BE} = 1 + 0,7 = 1,7 \text{ V d'où } R_2 = V_{R2} / I_2 = 1,7 / 0,01 = 170 \Omega$$

$$V_{R1} = E - V_{R2} = 10 - 1,7 = 8,3 \text{ V de même } R_1 = V_{R1} / I_1 = 8,3 / 0,01 = 830 \Omega$$

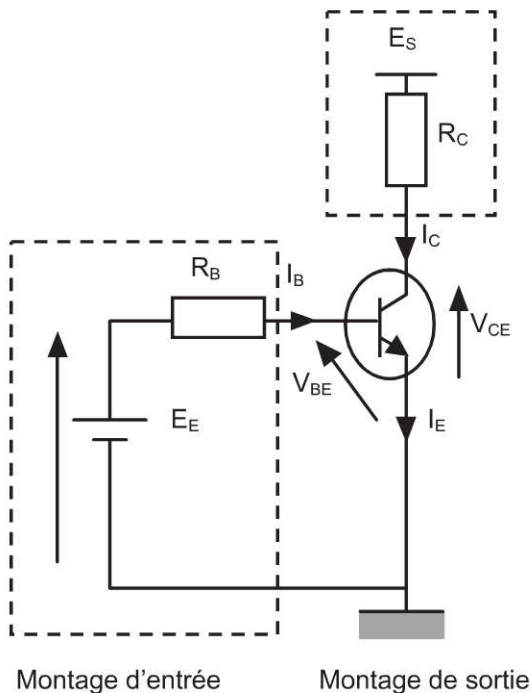
► Autre point de vue

Vu de la base, le schéma équivalent de montage en pont est identique à celui-ci.

On impose la valeur de la résistance d'entrée $R_1 // R_2 < R_{\text{entrée}}$ avec $R_{\text{entrée}} = \beta \times R_E$ pour satisfaire à la condition du vrai diviseur de tension entre R_1 et R_2 ($I_1 \cong I_2$)



b) Polarisation par la base d'un NPN



Il s'agit du même montage que celui utilisé pour la **fiche 33** « réseau de caractéristiques ».

Les relations sont :

$$\begin{cases} E_E = R_B \times I_B + V_{BE} \\ I_E = I_C + I_B = \beta \times I_B + I_B = (\beta + 1) \times I_B \cong I_C \\ V_{CE} = E - R_C \times I_C \end{cases}$$

Exemple

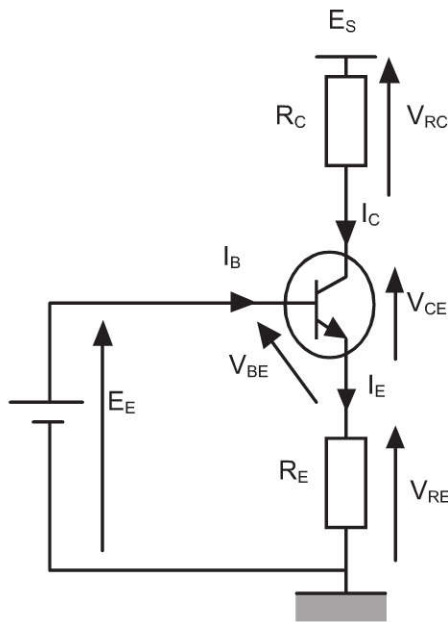
$$E_E = E_S = 15 \text{ V}, R_B = 500 \text{ k}, R_C = 3 \text{ k}, \beta_{\min} = 100.$$

$$I_B = 15 / 500 \cdot 10^3 = 30 \mu\text{A}$$

$$I_C = \beta \times I_B = 100 \times 30 \cdot 10^{-6} = 3 \text{ mA}$$

$$V_{CE} = 15 - 3 \cdot 10^3 \times 30 \cdot 10^{-3} = 6 \text{ V}$$

c) Polarisation par l'émetteur d'un NPN



L'application de la tension sur la base se fait de façon directe : plus besoin de limiter le courant I_B car cela va se faire naturellement en fixant I_E par le biais de R_E . La valeur de E_E peut être fixée par la relation $E_E = 5 \times V_{BE}$.

Ce montage rend le point de repos insensible aux variations de β donné par le constructeur pour une référence de transistor ($\beta_{\min}$ à $\beta_{\max}$). Cela se vérifie par les relations suivantes qui ne font pas appel à β pour obtenir la droite de charge.

$$\begin{cases} E_S = R_C \times I_C + V_{CE} + V_{RE} \\ I_E = V_{RE} / R_E \\ V_{RE} = E_E - V_{BE} \end{cases}$$

► Méthode applicative pour obtenir les valeurs des résistances

On impose $E_S = 10 \text{ V}$, $I_E \cong I_C = 5 \text{ mA}$, pour le transistor choisi $\beta_{\min} = 100$, $V_{BE} = 0,7 \text{ V}$.

On se fixe 50 % de E_S pour V_{CE} à l'état bloqué soit 5 V.

D'où $E_E = 5 \times 0,7 = 3,5 \text{ V}$ ce qui donne $V_{RE} = E_E - V_{BE} = 3,5 - 0,7 = 2,8 \text{ V}$.

$$R_E = V_{RE} / I_E = 2,8 / 5 \cdot 10^{-3} = 560 \Omega$$

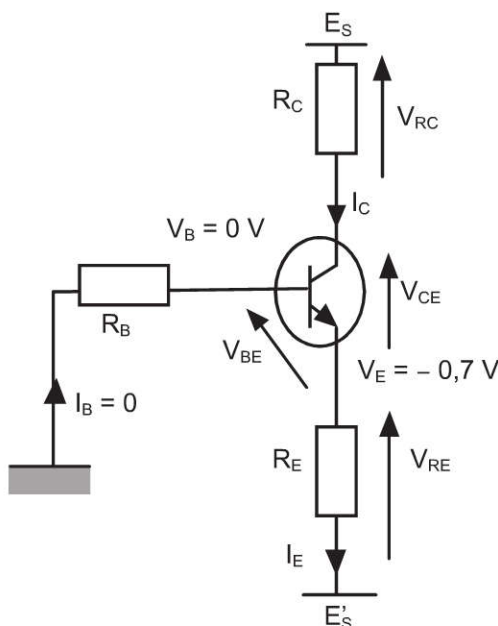
$$R_C \times I_C = E_S - V_{CE} - V_{RE} = 10 - 5 - 2,8 = 2,2 \text{ V} \text{ soit } R_C = 2,2 / 5 \cdot 10^{-3} = 440 \Omega$$

► Influence du gain β

L'influence du gain en courant n'est absolument pas significative puisque :

$$I_E = I_C + I_B = I_B + \beta \times I_B \quad I_B = I_E \times \frac{\beta}{\beta + 1} \cong I_E \text{ puisque } \beta \gg 1$$

d) Polarisation par l'émetteur d'un NPN avec deux sources



Le courant de base reste négligeable devant le I_C et I_E ($\beta \gg 1$), cela implique que la tension appliquée sur la base vaut 0V.

Puisque $V_B = 0 \text{ V}$ et $V_E = 0,6$ à $0,7 \text{ V}$ (tension aux bornes de la jonction base-émetteur) alors $V_{BE} = -0,6$ à $0,7 \text{ V}$.

Les relations de ce montage sont :

$$\begin{cases} E_S - (-E'_S) = R_C \times I_C + V_{CE} + V_{RE} \\ V_{RE} = -V_{BE} - (-E'_S) \\ V_{CE} = R_C \times I_C + V_{BE} \\ V_B \cong 0 \text{ V} \end{cases}$$

► **Méthode applicative pour obtenir les valeurs des résistances**

On impose $E_S = 10\text{ V}$, $E'_S = 5\text{ V}$, $I_E \cong I_C = 10\text{ mA}$ puisque R_B sera dimensionnée de façon à obtenir $I_B = 0$, pour le transistor choisi $\beta_{\min} = 150$, $V_{BE} = 0,7\text{ V}$. On se fixe 40 % de E_S soit $V_{CE} = 4\text{ V}$ à l'état bloqué.

$$V_{RE} = -0,7 - (-5) = 4,3\text{ V d'où } R_E = V_{RE} / I_E = 4,3 / 0,01 = 430\ \Omega$$

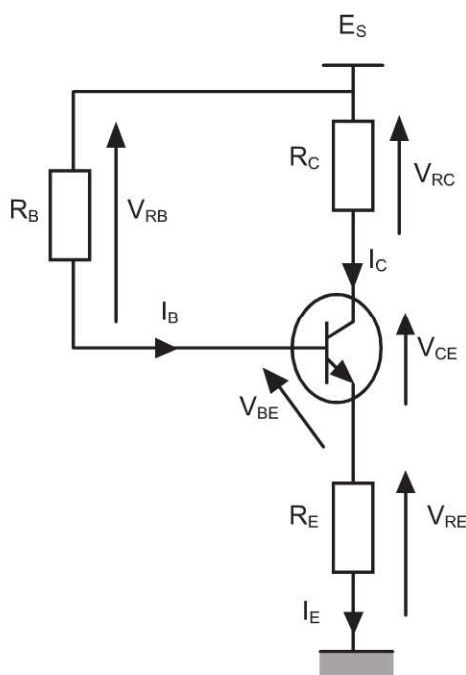
$$R_C \times I_C = E_S - V_{CE} - V_{RE} = 10 - 4 - 4,3 = 1,7\text{ V d'où } R_C = 1,7 / 0,01 = 170\ \Omega$$

Nous devons dimensionner R_B de façon à ce que I_B soit négligeable devant I_E .

$$I_B = I_C / \beta = 10 \cdot 10^{-3} / 150 = 66\ \mu\text{A avec un facteur 10 pour obtenir la condition souhaitée.}$$

$$I_B = 6,6\ \mu\text{A soit } R_B = 4,3 / 6,6 \cdot 10^{-6} = 700\ \text{k}\Omega.$$

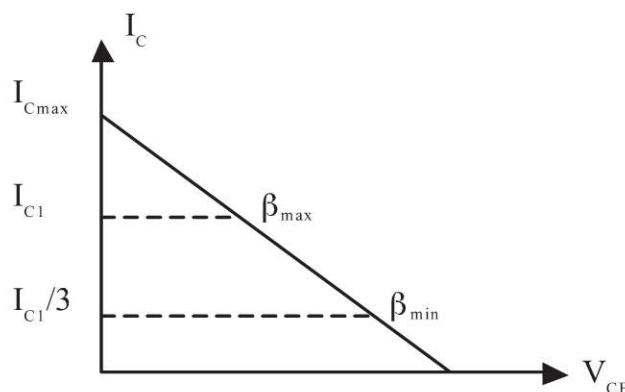
e) Polarisation par contre réaction de l'émetteur pour un NPN



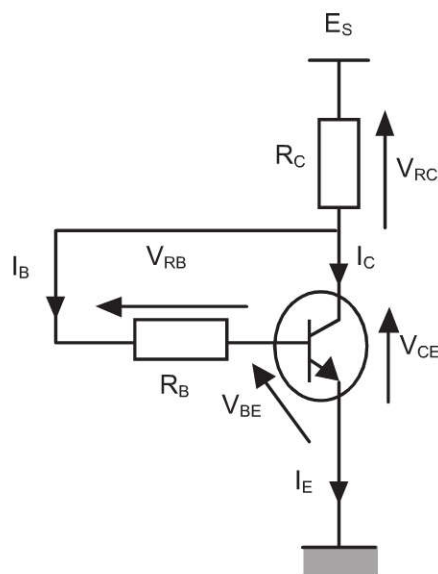
Le principe est le suivant : si I_C augmente, V_{RE} augmente, puisque $V_{BE} = 0,7\text{ V} = \text{constante}$, V_{RB} diminue, ce qui implique I_B qui diminue et donc un I_C qui diminue ($I_C = \beta \times I_B$).

Ce type de montage permet de palier aux variations en théorie. En pratique, il est très peu utilisé car $R_E > R_B / \beta$, ce qui implique une saturation ou un blocage très rapide du transistor. Dans le cas où $\beta_{\min}$ vaut 100 et $\beta_{\max}$ vaut 300, alors I_C sera modifié dans un rapport de 3.

$$\begin{cases} E_S = R_C \times I_C + V_{CE} + V_{RE} \\ V_{RE} = R_E \times I_E \\ V_{CE} = R_C \times I_C + V_{BE} \\ V_{RB} = E_S - V_{BE} - V_{RE} \\ V_{RE} + R_B \times I_B = E_S - V_{BE} \end{cases}$$



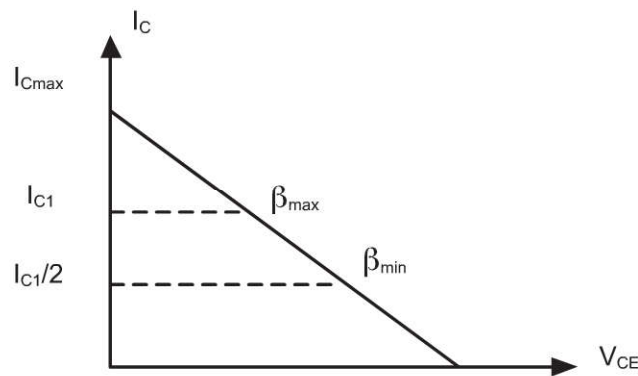
f) Polarisation par contre réaction du collecteur pour un NPN



Le principe reste similaire au précédent. Une augmentation de I_C entraîne une augmentation de V_{RC} donc une diminution de la tension V_{CE} et une diminution de V_{RB} . Ceci a pour effet de diminuer I_B et donc de s'opposer à l'augmentation de I_C .

Dans le cas où $\beta_{\min}$ vaut 100 et $\beta_{\max}$ vaut 300 alors I_C sera modifié dans un rapport de 2, ce qui est un peu mieux que la polarisation par contre réaction de l'émetteur.

$$\begin{cases} E_S = R_C \times I_C + V_{CE} \\ E_S - V_{BE} = \frac{I_E}{\beta} \times R_B \end{cases}$$



3. EN PRATIQUE

	Inconvénients	Avantage	Utilisation
Polarisation en pont	utilise quatre résistances	utilise qu'une source de tension et caractéristique de charge indépendante du β	amplification
Polarisation par la base	utilise deux sources de tensions, caractéristique de charge dépendante du β	maîtrise complète de la valeur du courant I_B	commutation
Polarisation par l'émetteur	utilise deux sources de tensions	utilisation de deux résistances, caractéristique de charge indépendante du β	amplification
Polarisation par l'émetteur à deux sources	utilise deux sources de tensions dont une de valeur négative	utilisation de deux résistances, caractéristique de charge indépendante du β	amplification
Polarisation par contre réaction de l'émetteur	caractéristique de charge très dépendante du β , utilisation de trois résistances	une seule source de tension	très peu utilisé
Polarisation par contre réaction de l'émetteur ou collecteur	caractéristique de charge très dépendante du β	utilisation de deux résistances, et une seule source de tension	amplification

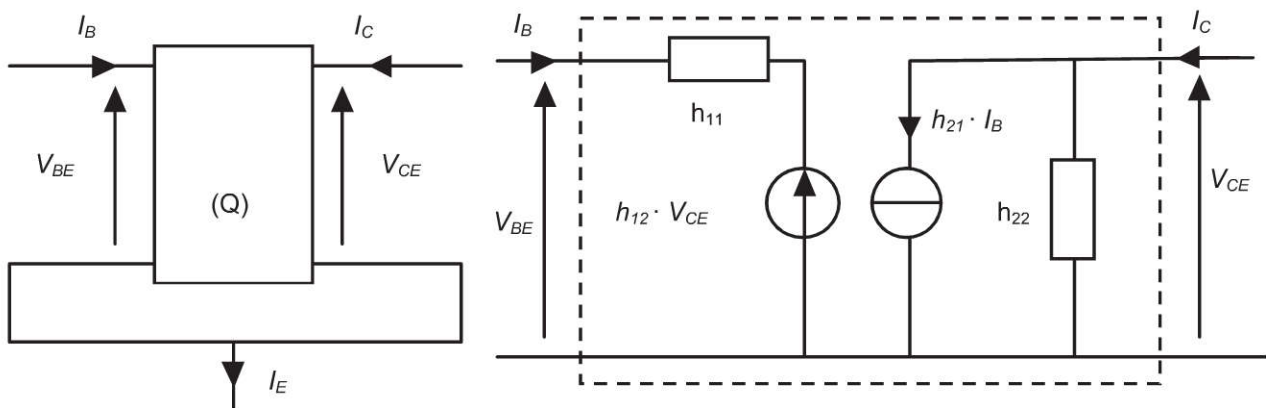
35 Modèle du transistor bipolaire

Mots clés

Paramètres, modèle, impédance, facteur d'amplification, conductance.

1. EN QUELQUES MOTS

Le transistor peut être vu comme un quadripôle Q dont les entrées seraient V_{BE} et I_B et les sorties V_{CE} et I_C , ce qui donne la représentation suivante dans le cas d'un émetteur commun. Il est important de signaler que les grandeurs I_B , I_C , V_{BE} , V_{CE} sont **ici des grandeurs variables sinusoïdales**.



2. CE QU'IL FAUT RETENIR

Le modèle proposé s'applique pour des **signaux sinusoïdaux évoluant autour d'un point de repos** fixé par un montage de polarisation. L'amplitude de ces signaux est faible devant les grandeurs de repos. Ceci permet de rester dans les zones linéaires de fonctionnement des différentes caractéristiques qui expliquent les valeurs élevées ou très faibles des paramètres. Le modèle s'appuie sur les caractéristiques d'entrée, de sortie, de transfert en courant (entrée – sortie), définies dans la **fiche 33** « réseau de caractéristiques ».

a) Les paramètres du transistor bipolaire

► Représentation hybrides des paramètres

$$\begin{cases} V_{BE} = h_{11} \cdot I_B + h_{12} \cdot V_{CE} \\ I_C = h_{21} \cdot I_B + h_{22} \cdot V_{CE} \end{cases}$$

Le paramètre h_{11} est une résistance, h_{12} et h_{21} sont sans dimensions, h_{22} est une conductance. L'étude du transistor bipolaire amène à négliger certains de ces paramètres.

► Détermination graphiques des paramètres

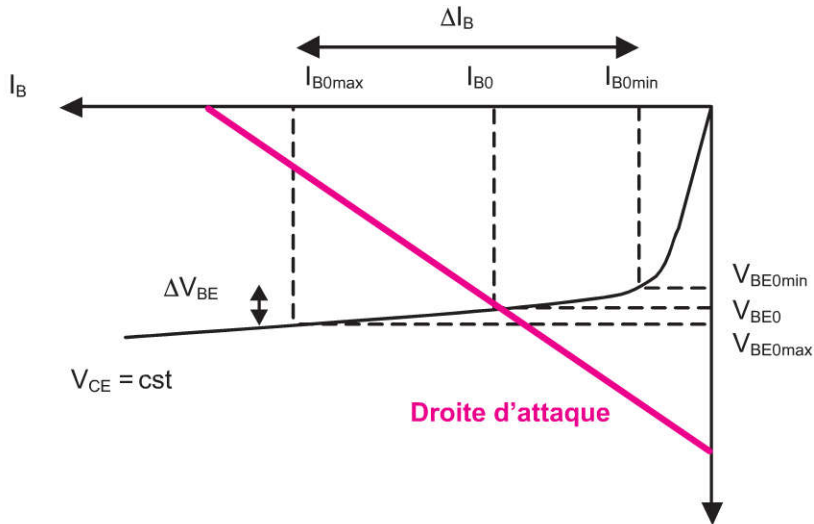
$$\begin{cases} \Delta V_{BE} = h_{11} \cdot \Delta I_B + h_{12} \cdot \Delta V_{CE} \\ \Delta I_C = h_{21} \cdot \Delta I_B + h_{22} \cdot \Delta V_{CE} \end{cases}$$

On note Δ l'indice signifiant qu'il s'agit d'une variation autour du point de repos fixé par le montage polarisant. Pour avoir un ordre d'idée des grandeurs, nous utiliserons les caractéristiques du transistor 2N3903.

b) Le paramètre h_{11} du transistor bipolaire (impédance d'entrée)

Le paramètre h_{11} correspond à la résistance différentielle de la diode d'entrée.

$$h_{11} = \left(\frac{\Delta V_{BE}}{\Delta I_B} \right)_{\Delta V_{CE}=0} = \left(\frac{\Delta V_{BE}}{\Delta I_B} \right)_{V_{CE}=cst}$$



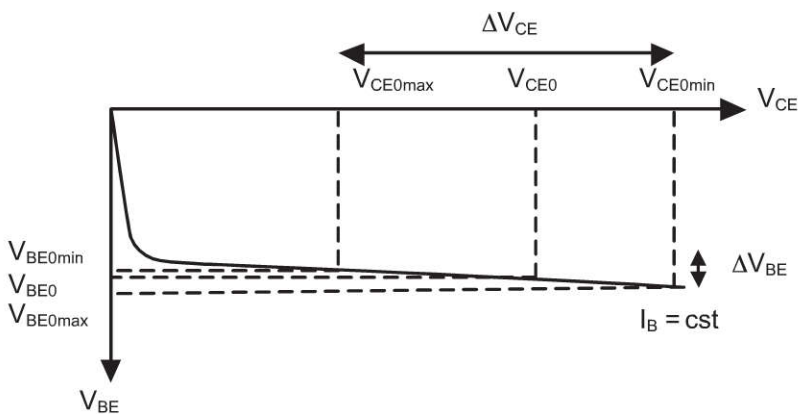
nom donné par la fiche constructeur $h_{11} = h_{ie}$.

ordre de grandeur min-max = de 1 kΩ à 8 kΩ

c) Le paramètre h_{12} du transistor bipolaire (facteur de rétroaction en tension)

Le paramètre h_{12} ne correspond à aucune grandeur physique : il est sans unité, il traduit juste le taux de réaction de la sortie sur l'entrée. La sortie est ouverte pour la partie alternative du signal.

$$h_{12} = \left(\frac{\Delta V_{BE}}{\Delta V_{CE}} \right)_{\Delta I_B=0} = \left(\frac{\Delta V_{BE}}{\Delta V_{CE}} \right)_{I_B=cst}$$



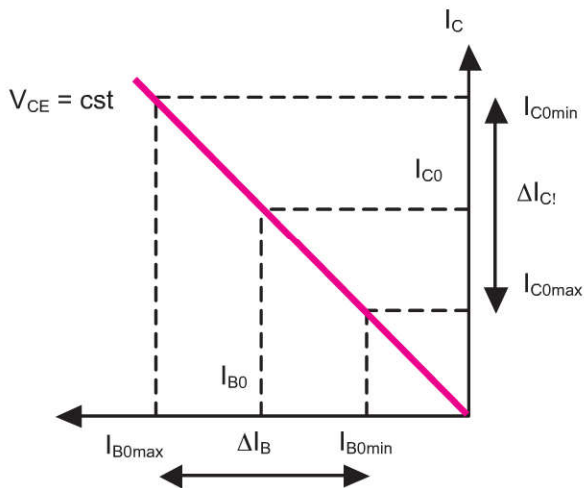
nom constructeur $h_{12} = h_{re}$.

ordre de grandeur min-max = de $0,1 \cdot 10^{-4}$ à $5 \cdot 10^{-4}$. La caractéristique étant quasi horizontale, la pente est extrêmement faible, le paramètre h_{12} est considéré comme nul.

d) Le paramètre h_{21} du transistor bipolaire (facteur d'amplification en courant)

Le paramètre h_{21} ne correspond à aucune grandeur physique, il est sans unité, il traduit l'amplification en courant entre l'entrée et la sortie. La sortie est en court-circuit pour la partie alternative du signal.

$$h_{21} = \left(\frac{\Delta I_C}{\Delta I_B} \right)_{\Delta V_{CE}=0} = \left(\frac{\Delta I_C}{\Delta I_B} \right)_{V_{CE}=cst}$$



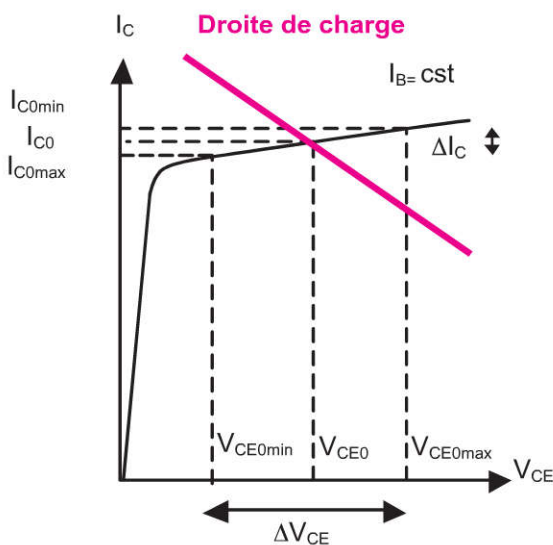
nom constructeur $h_{21} = h_{fe}$.

ordre de grandeur min-max = de 50 à 200

e) Le paramètre h_{22} du transistor bipolaire (conductance de sortie)

Le paramètre h_{22} correspond à l'inverse de la résistance de la sortie lorsque l'entrée est ouverte pour la partie alternative.

$$h_{22} = \left(\frac{\Delta I_C}{\Delta V_{CE}} \right)_{\Delta I_B = 0} = \left(\frac{\Delta I_C}{\Delta V_{CE}} \right)_{I_B = \text{cst}}$$



nom constructeur $h_{22} = h_{oe}$.

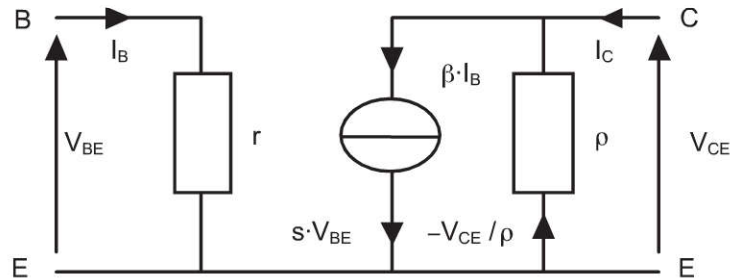
ordre de grandeur min-max = de 1 à 40 $\mu\text{Siemens}$
La pente est très faible ce qui implique une valeur, en siemens, de h_{12} faible (μS) et donc une résistance $1/h_{12}$ importante (M).

f) Les paramètres simplifiés et modèle équivalent du transistor bipolaire

Il est souvent plus pratique d'utiliser une version simplifiée des paramètres hybrides, ainsi nous définissons par r , β et ρ le modèle équivalent du transistor.

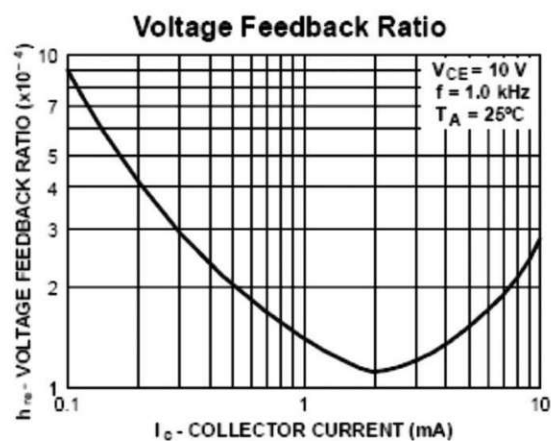
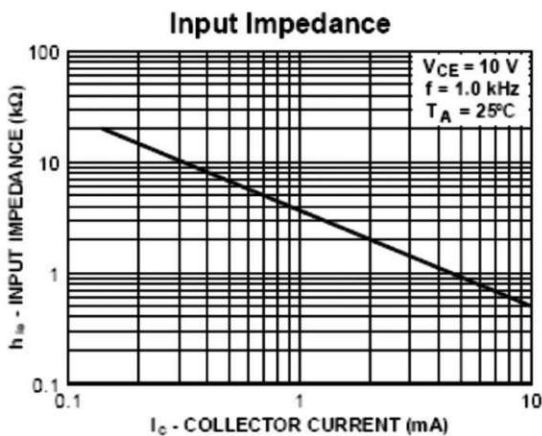
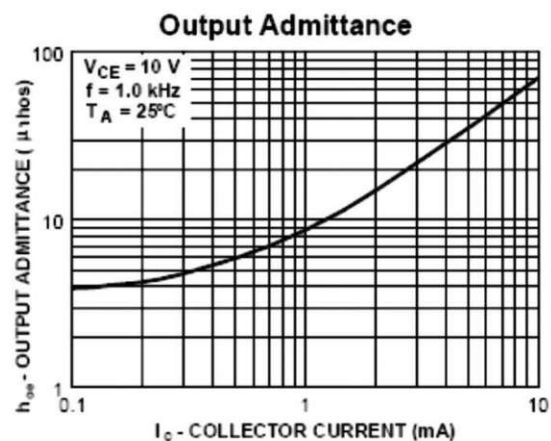
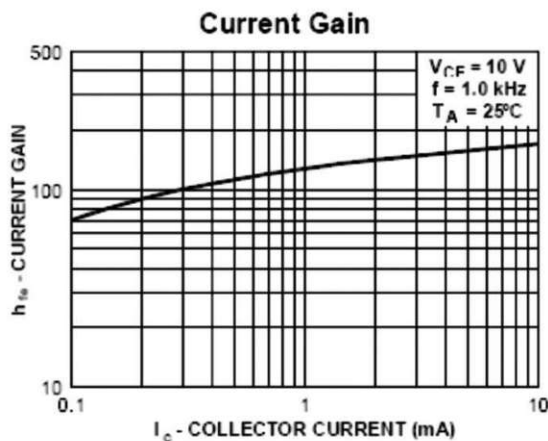
Une autre grandeur est introduite par les constructeurs qui est « s » nommée pente du transistor et exprimée en siemens. Dans les deux cas, $s \cdot V_{BE} = \beta \cdot I_B$ et le modèle reste rigoureusement identique.

$$\begin{aligned} h_{11} &= h_{ie} = r \\ h_{12} &= h_{re} = 0 \\ h_{21} &= h_{fe} = \beta \\ h_{22} &= h_{oe} = \frac{1}{\rho} \\ s &= \frac{h_{21}}{h_{11}} = \frac{\beta}{r} \end{aligned}$$

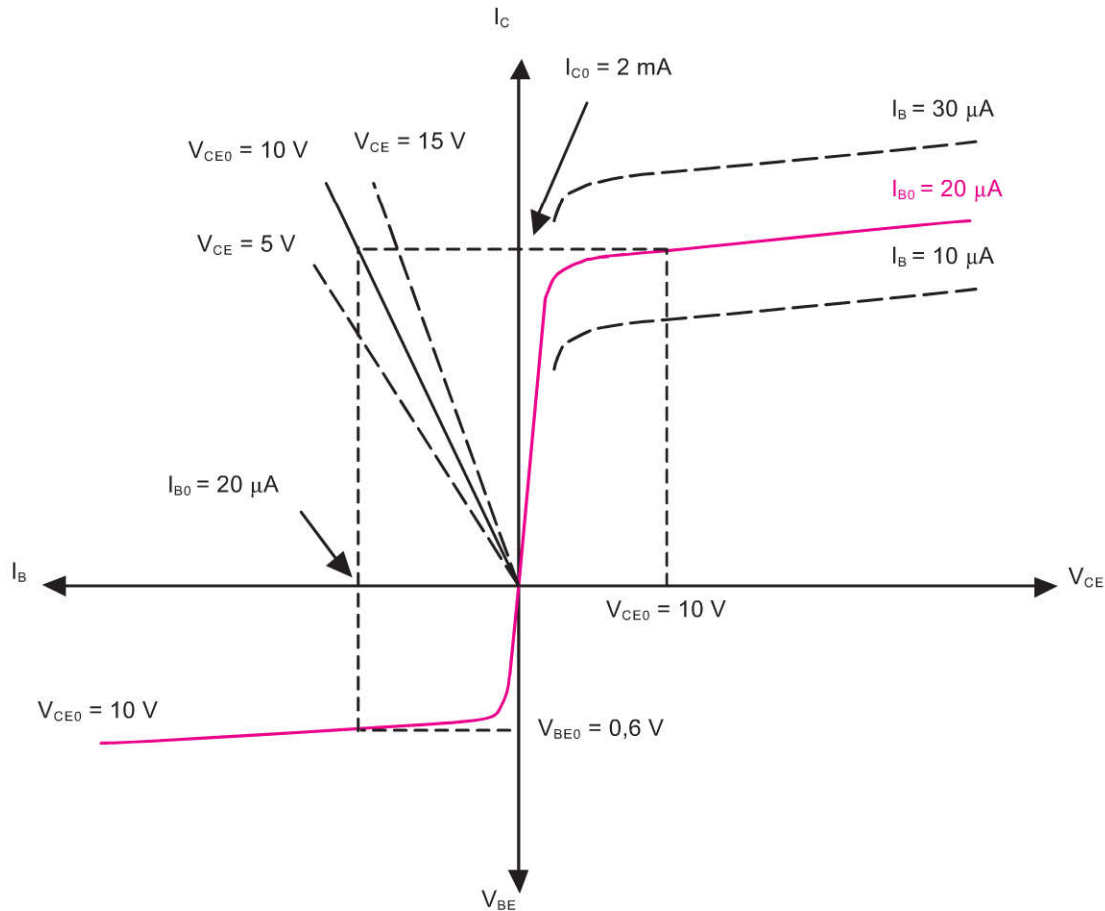


g) Variations des paramètres du modèle

Les paramètres du modèle sont soumis à variations, en fonction de la température mais notamment en fonction de I_C . Ce sont les variations qui sont indiquées dans les fiches constructeurs. Les caractéristiques ci-dessous montrent ces variations à température, fréquence des signaux (autour du point de repos), et à V_{CE} donnés.



Les différents quadrants peuvent être assemblés dans un seul et unique repère :



3. EN PRATIQUE

En utilisant le réseau de caractéristiques suivant, on peut déterminer les paramètres r , β et ρ au voisinage du point de repos ($V_{CE0} = 10 \text{ V}$ et $I_{B0} = 20 \mu\text{A}$) défini par un montage de polarisation.

Le signal alternatif a une amplitude qui permet de rester dans la zone linéaire de chaque caractéristique. Superposé au point de repos, il donne les variations suivantes :

$\Delta I_B = 10 \mu\text{A}$ soit une valeur $I_{B\text{max}} = 25 \mu\text{A}$ et $I_{B\text{min}} = 15 \mu\text{A}$.

$\Delta V_{BE} = 50 \text{ mV}$ soit une valeur $V_{BE\text{max}} = 650 \text{ mV}$ et $V_{BE\text{min}} = 550 \text{ mV}$.

$\Delta I_C = 1 \text{ mA}$ soit une valeur $I_{C\text{max}} = 2,5 \text{ mA}$ et $I_{C\text{min}} = 1,5 \text{ mA}$.

$\Delta V_{CE} = 16 \text{ V}$ soit une valeur $V_{CE\text{max}} = 18 \text{ V}$ et $V_{CE\text{min}} = 2 \text{ V}$.

$$r = \Delta V_{BE} / \Delta I_B = 0,05 / 10 \cdot 10^{-6} = 5 \text{ k}\Omega$$

$$\beta = \Delta I_C / \Delta I_B = 1 \cdot 10^{-3} / 10 \cdot 10^{-6} = 100$$

$$\rho = \Delta V_{CE} / \Delta I_C = 16 / 1 \cdot 10^{-3} = 16 \text{ k}\Omega.$$

a) Conclusion

Le modèle équivalent du transistor pour de petits signaux évoluant autour d'un point de repos est constitué de r (impédance d'entrée), β (facteur d'amplification en courant) et de ρ (conductance de sortie). Les paramètres sont donnés par la fiche constructeur où ils peuvent être déterminés graphiquement. Les grandeurs sont fonctions de la température et du courant I_C :

• **vu de l'entrée** : $V_{BE} = r \times I_B$

• **vu de la sortie** : $\beta \times I_B = I_C - \frac{V_{CE}}{\rho}$

36 Amplificateur à transistor (principe)

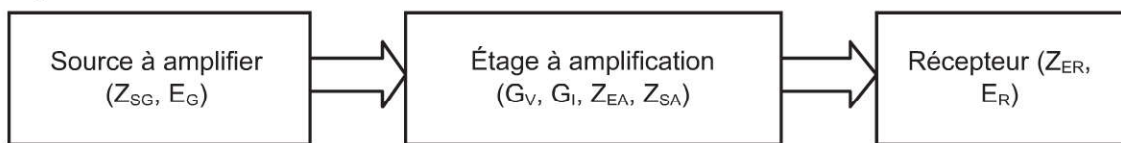
Mots clés

Polarisation, signal, alternatif, découpler, amplification.

1. EN QUELQUES MOTS

Le transistor étant un composant permettant de commander un courant avec un autre courant. Il peut être utilisé comme amplificateur de signaux. Pour cela, il faut que le transistor :

- soit polarisé, c'est-à-dire que l'on fixe le point de fonctionnement moyen des grandeurs d'entrées (V_{BE} et I_B) et de sortie (V_{CE} et I_C) ;
- le signal d'entrée ait une amplitude maximale permettant d'éviter la saturation des grandeurs de sorties (exemple : $V_{CE} >$ à l'alimentation, ce qui n'est pas possible) ;
- soit découplé des signaux continus, nécessaires à la polarisation, et des signaux alternatifs à amplifier.



E_G : tension efficace de la source d'entrée, E_R : efficace appliquée à la charge après amplification

Z_{SG} : impédance de la source vue par l'amplificateur, Z_{SA} : vu par le récepteur

Z_{EA} : impédance d'entrée de l'amplificateur, Z_{ER} : impédance d'entrée du récepteur.

G_V, G_I : gain en tension et en courant généré par l'amplificateur.

2. CE QU'IL FAUT RETENIR

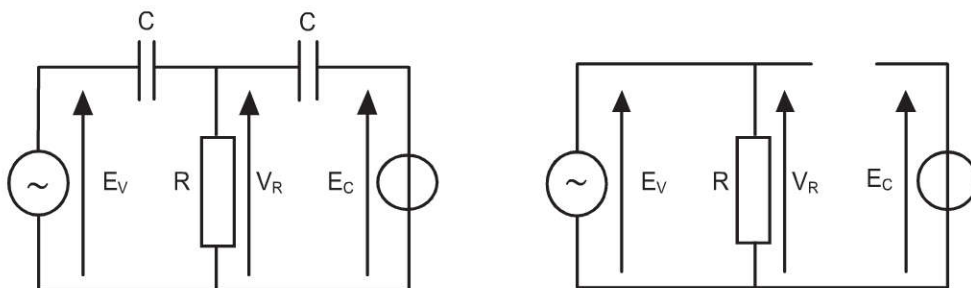
a) Séparation des signaux et capacités de découplage

Pour obtenir un découplage des signaux continus et alternatifs, il est nécessaire d'utiliser des condensateurs dits de découplage.

L'impédance d'un condensateur vaut $Z_C = 1 / C \cdot \omega$. Pour un signal de pulsation $\omega = 2 \cdot \pi \cdot f$ que l'on fait tendre vers l'infini, alors $Z_C = 0 \Omega$. Pour un signal continu, f tend vers 0 donc $Z_C = \infty$

Synthèse

E_C est une tension continue et E_V est la valeur efficace d'une tension alternative sinusoïdale. Pour f , le schéma à gauche est équivalent à celui de droite. V_R vaut alors E_V .



Quelle est la valeur de C à donner lorsque l'on connaît l'excursion en fréquence du signal à amplifier ? On peut appliquer la règle suivante : il faut que $Z_C < 0,1 \cdot R$ c'est-à-dire que l'impédance du

condensateur soit 10 fois plus faible que la résistance R à laquelle on veut transmettre le signal alternatif.

Exemple

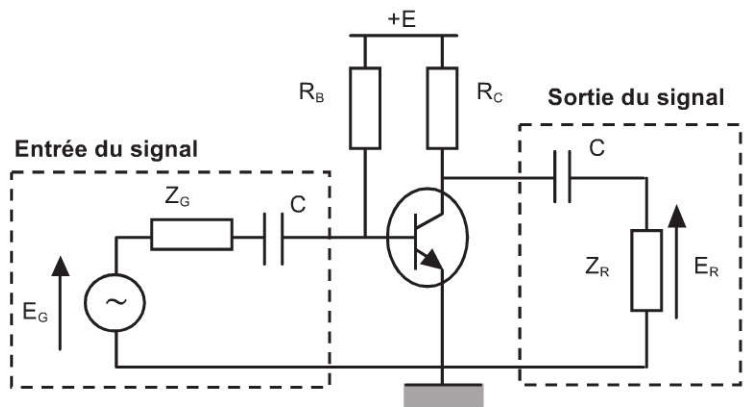
La tension $E_G(t)$ peut varier de 10 kHz à 100 kHz, pour une résistance d'entrée du montage de 50 Ω alors il faut que $Z_C < 5$ pour une fréquence minimale de 10 kHz. Soit $C > 3,1 \mu\text{F}$.

b) Les amplificateurs à transistor bipolaire : schéma de principe pour les petits signaux

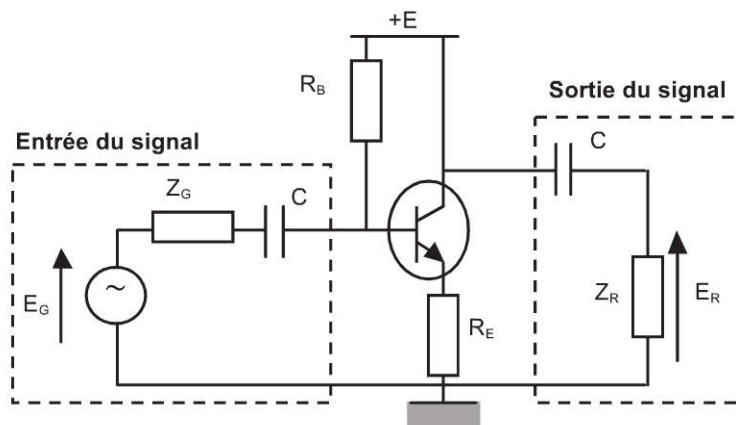
Les schémas de principes suivants permettent de repérer l'entrée du signal à amplifier et la sortie du signal amplifié. Nous faisons abstraction du montage polarisant nécessaire à la mise en place du point de repos. Le montage de polarisation est proposé par principe mais n'est pas exclusif.

Émetteur commun (fiche 37)

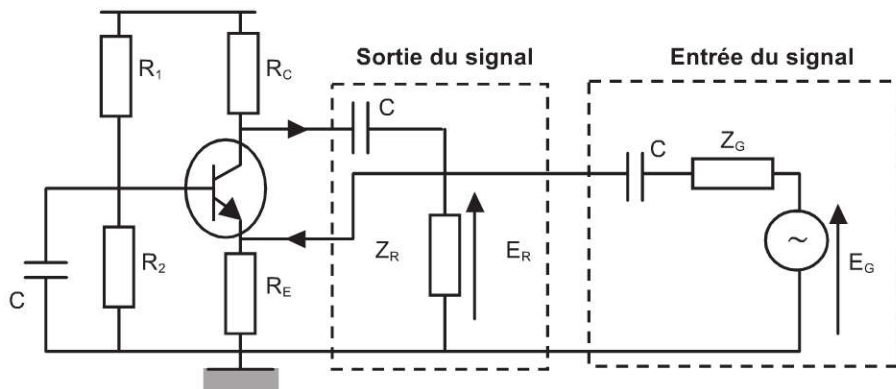
Le signal découplé est émis depuis la base et transmis vers le collecteur, l'émetteur étant mis à la masse du montage de polarisation. Une variante permet de faire intervenir une résistance sur l'émetteur pour la polarisation qui est découplée par un condensateur de liaison en parallèle. Cette dernière permet de modifier certaines caractéristiques du montage.



Collecteur commun (fiche 38)



Base commune (fiche 39)



37 Amplificateur émetteur commun

Mots clés

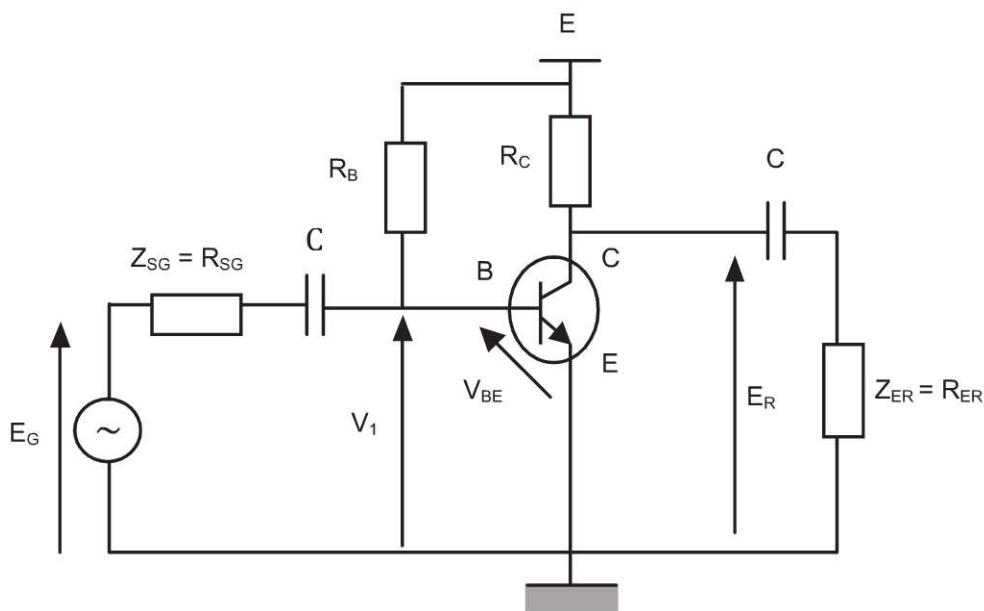
Gain, impédance d'entrée et de sortie, influence, récepteur.

1. EN QUELQUES MOTS

Le montage à transistor bipolaire émetteur commun offre une grande amplification en tension et en courant avec une impédance d'entrée et une impédance de sortie relativement faibles. Il est souvent associé à des étages adaptateurs d'impédance.

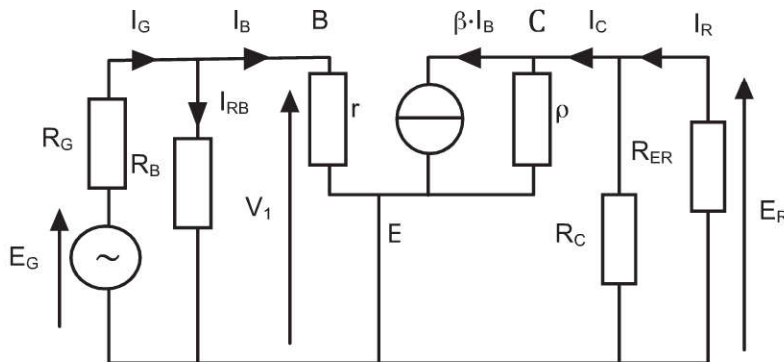
2. CE QU'IL FAUT RETENIR

On suppose l'impédance de sortie du signal à amplifier purement résistive R_{SG} et l'impédance du récepteur, elle aussi purement résistive (R_{ER}), la polarisation se faisant par les résistances R_B et R_C et la source E . Les condensateurs de liaisons C sont dimensionnés pour être considérés comme des courts-circuits à la fréquence du signal E_G . Le générateur E n'intervient pas pour les tensions alternatives, il est donc considéré comme un court-circuit entre E et la masse.



a) Schéma équivalent pour les petits signaux

On pose R_{Br} la résistance équivalente $R_B // r$; de même $R_T = \rho // R_C // R_{ER}$. On rappelle que toutes les grandeurs courant et tension sont sinusoïdales. Équations déduites du schéma :



$$R_{pRC} = \frac{1}{\frac{1}{\rho} + \frac{1}{R_C}} \quad (1) = \rho // R_C$$

$$I_B = \frac{R_B}{R_B + r} I_G \quad (2)$$

$$I_R = \frac{R_{pRC}}{R_{pRC} + R_{ER}} \beta \times I_B \quad (3)$$

$$R_T = \rho // R_C // R_{ER}$$

$$R_{Br} = R_B // r$$

b) Performances de ce montage

► Gain en courant G_I

En partant des relations (2) et (3) on obtient :

$$G_I = \frac{I_R}{I_G} = \beta \times \frac{R_B}{R_B + r} \times \frac{R_{pRC}}{R_{pRC} + R_{ER}}$$

► Gain en tension G_V

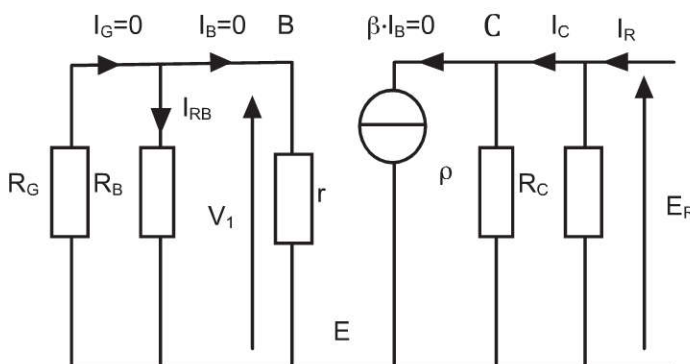
Le gain en tension peut être déduit de G_I , ainsi :

$$\begin{aligned} R_{Br} \times I_G &= V_1 \\ R_T \times I_R &= -E_R \end{aligned} \quad G_V = \frac{E_R}{V_1} = -\frac{R_{Br} \times I_R}{R_T \times I_G} = -\frac{R_{Br}}{R_T} \times G_I$$

Il y a opposition de phase entre la tension de sortie et la tension d'entrée.

• Impédance d'entrée Z_{EA} :

$$Z_{EA} = \frac{V_1}{I_G} = \frac{R_B \times r}{R_B + r} = R_{Br}$$

• Impédance de sortie Z_{SA} :

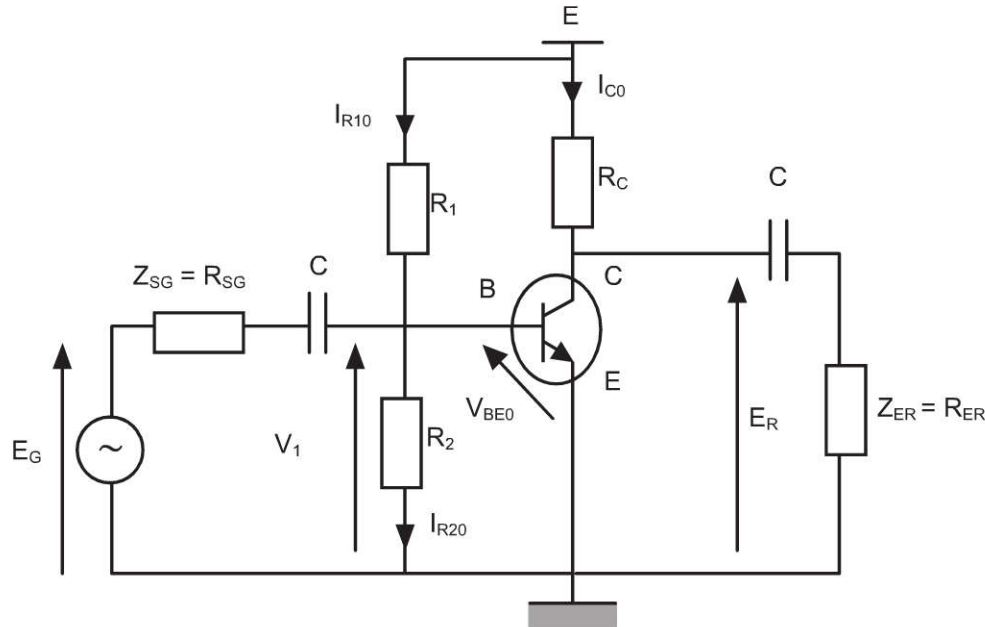
L'impédance de sortie est déduite lorsque le générateur de tension E_G d'entrée est un court-circuit et que la charge est déconnectée. $I_B = 0$ donc $\beta \times I_B = 0$

$$Z_{SA} = \frac{E_R}{I_R} = R_{pRC}$$

3. EN PRATIQUE

Le montage suivant est polarisé par les résistances R_1 , R_2 et R_C . On donne $\beta = 100$, ρ infini, $E = 12$ V, $r = 2$ k Ω , $\rho = \infty$, $R_C = 1$ k Ω , $R_1 = 20$ k Ω , $R_{ER} = 1$ k Ω , on désire un point de repos dont les valeurs sont $I_{C0} = 5$ mA, $V_{BE0} = 0,7$ V.

Les grandeurs permettant la mise en place du point de repos sont notées avec un indice 0.



a) Pour le point de repos

Détermination de R_2 :

$$I_{B0} = I_{C0} / \beta = 5 \text{ mA} / 100 = 50 \text{ } \mu\text{A}, V_{R20} = V_{BE0} = 0,7 \text{ V}, V_{R10} = E - V_{BE0} = 12 - 0,7 = 11,3 \text{ V}.$$

$$I_{R10} = V_{R10} / R_1 = 11,3 / 20 \cdot 10^3 = 0,565 \text{ mA}, I_{R20} = I_{R10} - I_{B0} = 0,565 - 0,05 = 0,56 \text{ mA}.$$

$$R_2 = V_{R20} / I_{R20} = 0,7 / 0,56 \text{ mA} = 1\,250 \text{ } \Omega.$$

b) Pour le régime dynamique

$$R_B = R_1 // R_2 = 1\,176 \text{ } \Omega$$

► Impédance d'entrée : $Z_{EA} = \frac{V_1}{I_G} = \frac{R_B \times r}{R_B + r} = \frac{1176 \times 2000}{1176 + 2000} = R_{Br} = 740 \text{ } \Omega$

► Impédance de sortie : $Z_{SA} = \frac{E_R}{I_R} = R_{pRC} = R_C = 1000 \text{ } \Omega$

► Gain en courant :

$$G_I = \frac{I_R}{I_G} = \beta \times \frac{R_B}{R_B + r} \times \frac{R_{pRC}}{R_{pRC} + R_{ER}} = 100 \times \frac{1176}{1176 + 2000} \times \frac{1000}{1000 + 1000} = 18,5$$

► Gain en tension :

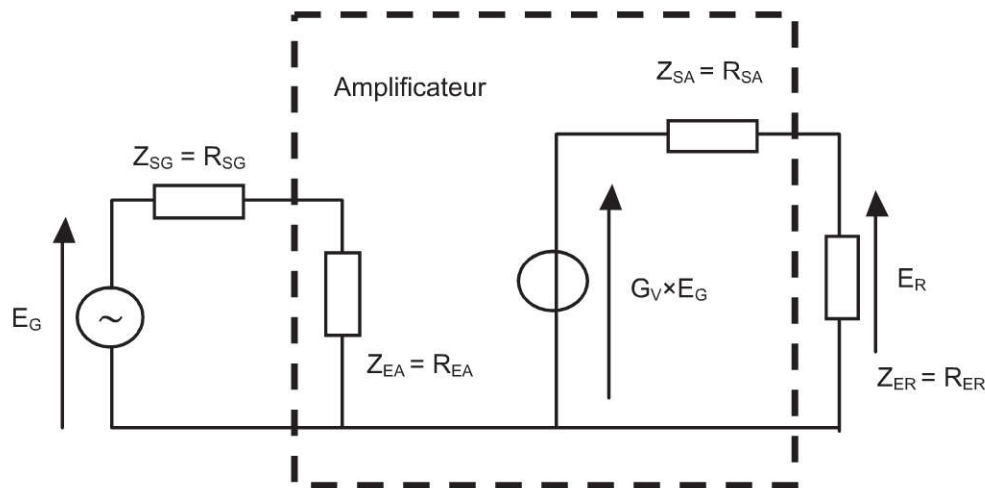
$$G_V = \frac{E_R}{V_1} = - \frac{R_{Br} \times I_R}{R_T \times I_G} = - \frac{R_{Br}}{R_T} \times G_I = - \frac{740}{500} \times 18,5 = -27,4$$

c) Conclusion

- **Avantage** : grande amplification en tension et en courant.
- **Inconvénients** : impédance d'entrée relativement faible (R_B/r), ce qui génère une grande consommation de puissance de la part du générateur devant délivrer le signal à amplifier. Impédance de sortie relativement petite d'où forte consommation du montage pour amplifier le signal. Déphasage de 180° entre les tensions d'entrée et de sortie.
- **Application** : ce montage est l'amplificateur de base à transistor et sera donc utilisé comme sous-fonction dans des circuits plus complexes (discrets, ou intégrés comme dans l'amplificateur opérationnel). Par contre, il sera souvent inexploitable seul, et il faudra lui adjoindre des étages adaptateurs d'impédance. (fiche 38)

L'idéal pour un amplificateur serait :

- d'absorber un courant I_G au générateur nul, ce qui implique une impédance d'entrée Z_{EA} qui serait infinie.
- de délivrer une tension E_R au récepteur qui soit l'image exacte de E_G mais amplifiée d'où $Z_{SA} = 0$.
- sans distorsion (caractéristique de sortie linéaire).
- une consommation propre nulle.
- une amplification constante quelle que soit la fréquence du signal à amplifier.



Tout ceci n'est pas réalisable en un seul montage, mais l'association de plusieurs montages en cascade permet de satisfaire un certain nombre de paramètres.

38 Montage collecteur commun

Mots clés

Gain en tension unitaire, impédance d'entrée grande, impédance de sortie faible.

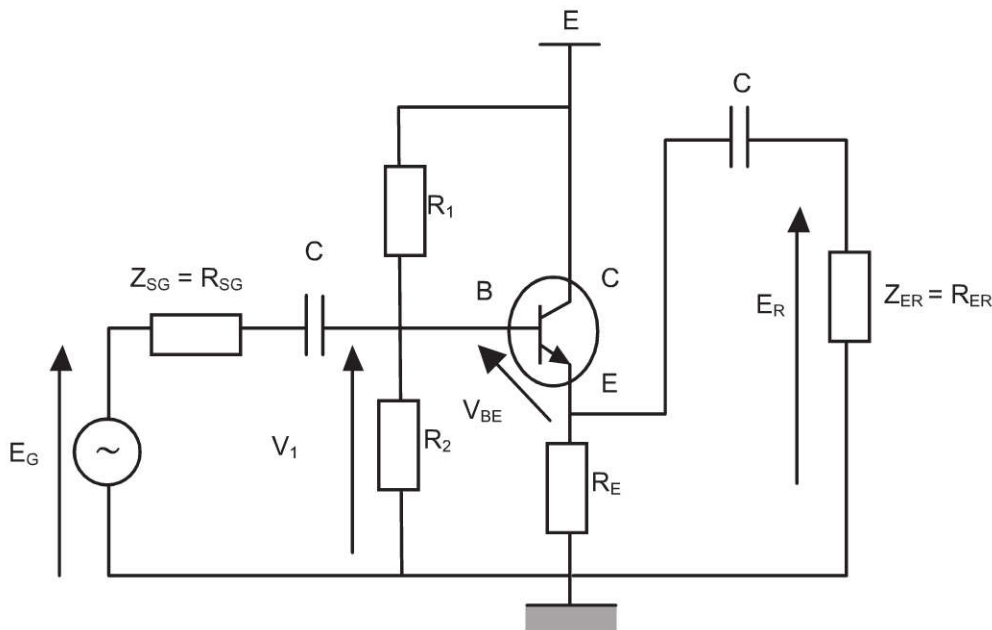
1. EN QUELQUES MOTS

Le montage à transistor bipolaire collecteur commun n'est pas un amplificateur, son gain étant proche de 1. Cependant, sa grande impédance d'entrée et sa faible impédance de sortie en font un système adaptateur d'impédance, particulièrement utilisé en association avec un amplificateur à émetteur commun.

2. CE QU'IL FAUT RETENIR

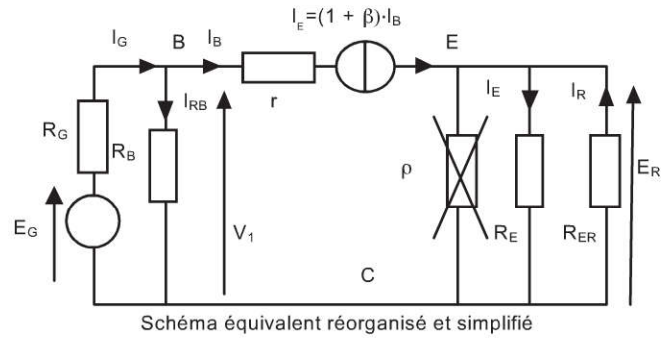
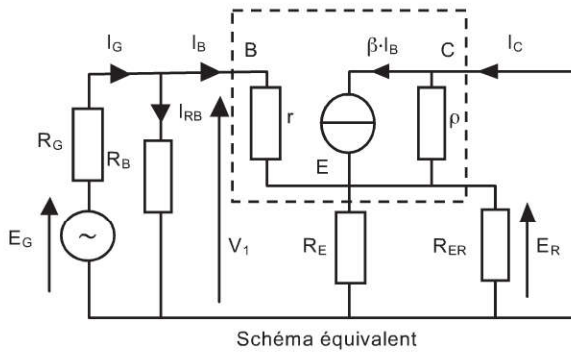
a) Montage

Les conditions sur le signal à amplifier sont les mêmes que pour l'amplificateur à émetteur commun. Pour la présentation de ce montage, nous utiliserons une polarisation avec les résistances R_1 , R_2 et R_E . Une fois de plus, ce mode de polarisation n'est pas exclusif à l'amplification à collecteur commun.



b) Schéma équivalent pour les petits signaux

On pose $R_B = R_1 // R_2$ et $R_T = R_E // R_{ER}$. Pour simplifier l'étude est souvent négligée devant R_E et R_{ER} . En effet, ρ est de l'ordre de la dizaine de $k\Omega$ ce qui est grand pour les valeurs courantes de R_E et R_{ER} . Cela signifie que $\rho // R_E // R_{ER} \cong R_E // R_{ER}$. On rappelle que toutes les grandeurs en courant et tension sont sinusoïdales.



c) Équations déduites du schéma

$$V_1 = R_B \times I_{RB} = r \times I_B + (1 + \beta) \times I_B \times R_T \quad E_R = -R_{ER} \times I_R = (\beta + 1) \cdot R_T \cdot I_B$$

$$R_T = \frac{1}{\frac{1}{R_E} + \frac{1}{R_{ER}}} \quad (2)$$

$$I_G = I_{RB} + I_B \quad (3)$$

$$I_R = \frac{-(\beta + 1) \times I_B \times R_E}{R_{ER} + R_E}$$

$$I_E = I_R + (\beta + 1) \times I_B = \frac{R_{ER}}{R_{ER} + R_E} \times (\beta + 1) \times I_B \quad (4)$$

$$(4) \Rightarrow I_R = (\beta + 1) \cdot I_G \cdot \frac{R_B}{R_S \times Z'_E} \left(1 - \frac{R_{ER}}{R_{ER} + R_E} \right)$$

On peut définir une impédance intermédiaire qui est notée Z'_E correspondant à l'impédance du montage vue aux bornes de R_B en partant de (1).

$$Z'_E = \frac{V_1}{I_B} = r + (\beta + 1) \times R_T$$

Dans ce cas, on peut définir $I_{BR} = \frac{Z'_E}{Z'_E + R_B} \times I_G$ en injectant dans l'expression (3) et on obtient :

$$I_B = \left(1 - \frac{Z'_E}{Z'_E + R_B} \right) \times I_G = -I_G \times \left(\frac{R_B}{R_B + Z'_E} \right) \quad (6)$$

On injecte (6) dans (4) et on obtient une expression qui est fonction de I_G et I_R .

d) Performances de ce montage

► Gain en courant G_I

$$I_R + (\beta + 1) \times \left(-I_G \times \frac{R_B}{R_B + Z'_E} \right) = \frac{R_{ER}}{R_{ER} + R_E} \times (\beta + 1) \times \left(-I_G \times \frac{R_B}{R_B + Z'_E} \right)$$

$$G_I = \frac{I_R}{I_G} = -(\beta + 1) \times \frac{R_B}{R_B + Z'_E} \times \frac{R_E}{R_E + R_{ER}}$$

Il y a opposition de phase entre le courant entrant I_G et le courant sortant I_R .

► Gain en tension G_V

$$G_V = \frac{E_R}{V_1} = \frac{(\beta + 1) \times R_T}{r + (\beta + 1) \times R_T} \cong 1$$

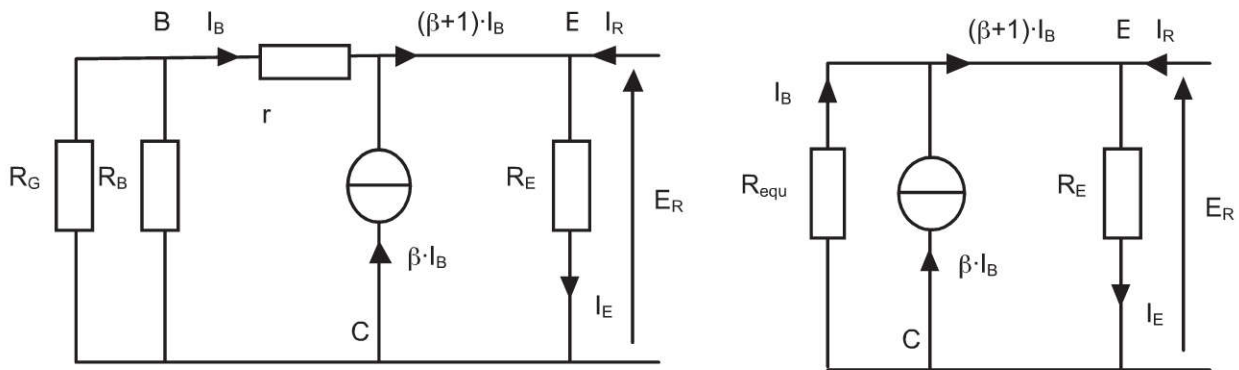
Le gain en tension est proche de l'unité car r est très inférieur $(\beta + 1) \times R_T$

► Impédance d'entrée Z_{EA}

$$Z_{EA} = \frac{V_1}{I_G} = \frac{R_B \times Z'_E}{R_B + Z'_E}$$

► Impédance de sortie Z_{SA}

L'impédance de sortie est déduite lorsque le générateur de tension E_G d'entrée est un court-circuit et que la charge est déconnectée. On applique une tension E_R et on veut définir le courant I_R . On note la résistance équivalente R_{EQU} la résistance formée de r en série avec $R_G // R_B$.



$$I_R = I_E - (\beta + 1) \times I_B = \frac{E_R}{R_E} - (\beta + 1) \times \frac{-E_R}{R_{EQU}}$$

d'où

$$Z_{SA} = \frac{E_R}{I_R} = \frac{R_E \times R_{EQU}}{R_{EQU} + R_E(\beta + 1)}$$

3. EN PRATIQUE

Le montage est celui donné en début de fiche. Les paramètres du transistor ont pour valeurs : $r = 1\,500\,\Omega$, $\beta = 200$, on suppose $\rho = \infty$. On impose $R_2 = 200\,\Omega$, $R_G = 500\,\Omega$, $R_E = R_{ER} = 1\,\text{k}\Omega$, $E = 10\,\text{V}$, $V_{CE0} = 5\,\text{V}$, $V_{BE0} = 0,7\,\text{V}$

Pour le point de repos on veut R_1 .

$$I_{C0} = (E - V_{CE0}) / R_E = 5 / 1\,000 = 5\,\text{mA}$$

$$I_{B0} = I_{C0} / \beta = 25\,\mu\text{A}, \quad V_{BE0} = V_{R20} = R_2 \times I_{20} \text{ d'où } I_{20} = 0,7 / 200 = 1,4\,\text{mA}.$$

$$I_{10} = I_{20} + I_B \quad I_{20} = 1,4\,\text{mA} \text{ d'où } R_1 = (E - V_{R20}) / I_{10} = (10 - 0,7) / 1,4 \cdot 10^{-3} = 6\,600\,\Omega$$

Pour les petits signaux

$$R_T = R_E // R_{ER} = 500, \quad R_B = R_1 // R_2 \quad R_1 = 6,6\,\text{k}\Omega$$

$$G_I = -6,1 \text{ avec } Z'_E = 102\,\text{k}\Omega$$

$$G_V \approx 1 \text{ et } Z_{EA} = Z'_E // R_B \text{ puisque } Z'_E \gg R_B \text{ alors } Z_{EA} \approx Z'_E$$

$$R_{EQU} = R_G // R_B + r \approx R_G + r \approx 2\,000\,\Omega \text{ ce qui donne } Z_{SA} = 500\,\Omega$$

a) Conclusion

Le montage à collecteur commun permet d'obtenir :

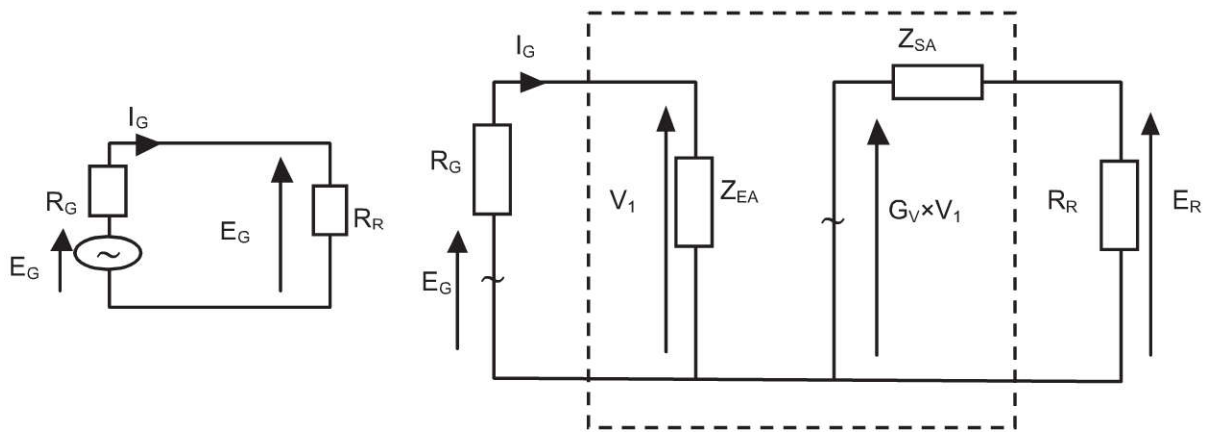
- **Avantages** : une grande impédance d'entrée et une petite impédance de sortie.
- **Inconvénients** : un gain en tension proche de l'unité quelle que soit la charge. Un gain en courant qui est fonction des résistances de polarisation (R_1 , R_2 , R_E) mais aussi de la charge.
- **Application** : ce montage est utilisé pour adapter l'impédance, il permet de faire le transfert du signal provenant d'une source vers une charge sans modifier le gain (fiche 21).

Exemple

On désire transférer le signal E_G aux bornes de R_R avec $R_G = 1\text{ k}\Omega$ et $R_R = 100\text{ }\Omega$

Par le diviseur de tension on obtient :

$$E_R = E_G \times \frac{R_R}{R_R + R_G} \cong \frac{E_G}{10}$$



Grâce au montage à collecteur commun, puisque $Z_{EA} \gg R_G$, $Z_{SA} \ll R_R$ et $G_V \cong 1$ alors

$$E_R = G_V \times V_1 \times \frac{R_R}{Z_{SA} + R_R} \cong E_G$$

$$V_1 = E_G \times \frac{Z_{EA}}{Z_{EA} + R_G} \cong E_G$$

Sans avoir à amplifier, on gagne un facteur 10 entre la sortie et l'entrée, perdu naturellement par la présence de R_G .

39 Amplificateur base commune

Mots clés

Gain en tension important, gain en courant unitaire, impédance d'entrée grande, impédance de sortie faible.

1. EN QUELQUES MOTS

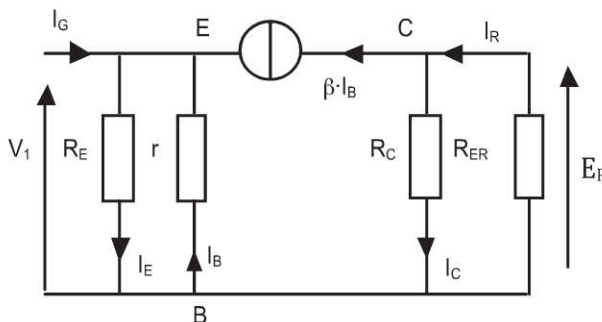
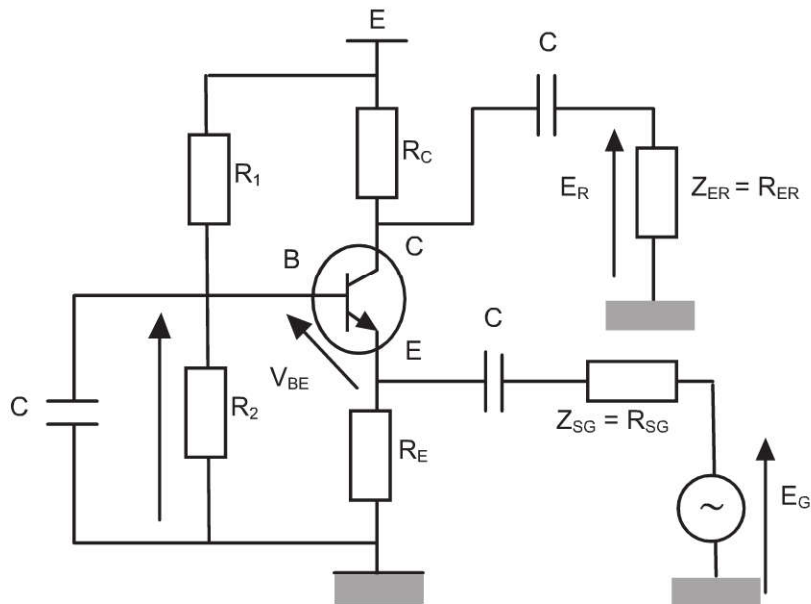
L'amplificateur à base commune est un amplificateur de tension sans inversion de phase dont l'utilisation est essentiellement destinée aux hautes fréquences.

2. CE QU'IL FAUT RETENIR

La base du transistor est court-circuitée par un condensateur de liaison pour le signal alternatif généré par E_G . Ce dernier se trouve donc connecté à la masse du montage. La polarisation se fera en pont grâce aux résistances R_1 , R_2 , R_C et R_E .

a) Schéma équivalent pour les petits signaux

On pose $R_T = R_C // R_{ER}$. Il est à noter que les résistances R_1 et R_2 ne figurent pas sur le schéma équivalent, simplement car chacune de leurs extrémités est reliée à la masse pour les petits signaux.



$$E_R = -r \times I_b(1)$$

$$E_R = -R_E \times I_E(2)$$

$$I_G = I_E - (\beta + 1) \times I_B(3)$$

$$E_R = -R_{ER} \times I_R(4) = -R_T \times \beta \times I_B(5)$$

$$I_R = \beta \times I_B \times \frac{R_C}{R_C + R_{ER}}(6)$$

En faisant (1) = (2) nous obtenons : $I_E = -\frac{r}{R_E} \times I_B$

puis en remplaçant dans (3), nous obtenons :

$$I_G = -\frac{r}{R_E} \times I_B - (\beta + 1) \times I_B = -\left(\frac{r}{R_E} + \beta + 1\right) \times I_B$$

b) Performances de ce montage

► Gain en courant G_I

En introduisant (5) et (6) dans la relation précédente, nous obtenons :

$$G_I = \frac{I_R}{I_G} = - \frac{\beta \times R_C}{\left(\frac{r}{R_E} + \beta + 1 \right) \times (R_C + R_{ER})}$$

Il y a opposition de phase entre le courant entrant I_G et le courant sortant I_R . Le gain G_I est < 1 .

► Gain en tension G_V

$$G_V = \frac{E_R}{V_1} = \frac{-R_{ER} \times I_R}{-r \times I_B} = \frac{1}{r \times I_B} \times R_{ER} \times \beta \times I_B \times \frac{R_C}{R_C + R_{ER}} = \beta \times \frac{R_T}{r}$$

Contrairement au montage à émetteur commun, le gain en tension est positif, ce qui signifie que les tensions E_G et E_R sont en phase.

► Impédance d'entrée Z_{EA}

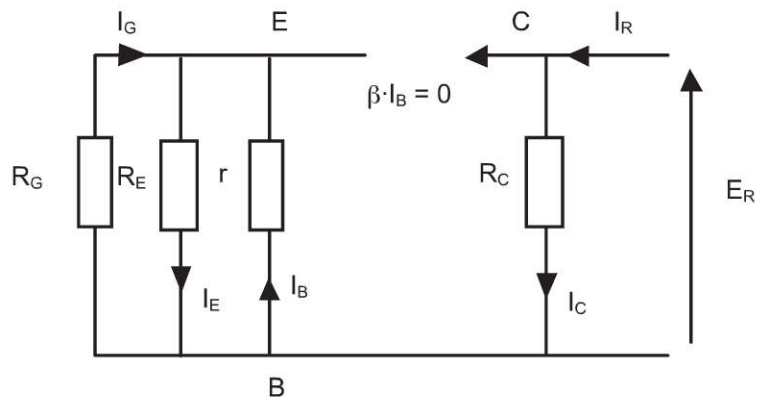
$$Z_{EA} = \frac{1}{Y_{EA}} = \frac{I_G}{E_G} = \frac{-\left(\frac{r}{R_E} + \beta + 1 \right) \times I_B}{-r \times I_B} = \frac{1}{R_E} + \frac{\beta + 1}{r} = \frac{r + (\beta + 1) \times R_E}{r \times R_E} \rightarrow Z_{EA} = \frac{r \times R_E}{r + (\beta + 1) \times R_E}$$

► Impédance de sortie Z_{SA}

L'impédance de sortie est déduite lorsque le générateur de tension E_G d'entrée est un court-circuit et que la charge est déconnectée. Du point de vue de la charge, il n'y a pas de calcul à effectuer, car l'impédance de la source de courant est infinie.

Donc on peut en déduire :

$$Z_{SA} = \frac{E_R}{I_R} = R_C$$



3. EN PRATIQUE

Pour le montage présenté en début de fiche, on pose $r = 500$, $\beta = 100$, $R_E = R_{ER} = 1 \text{ k}\Omega$, $R_C = 4 \text{ k}\Omega$. $G_I = (100 \times 4\,000) / ((0,5 + 100 + 1) \times 5\,000) = -0,79$. $G_V = (100 \times 4\,000 \times 1\,000) / (500 \times 5\,000) = 160$. $Z_{EA} = 4,9 \Omega$. $Z_{SA} = 4 \text{ k}\Omega$.

a) Conclusion

- **Avantages** : l'amplificateur à base commune n'amplifie pas en courant mais présente une amplification en tension importante et sans inversion de phase (contrairement à l'émetteur commun). Faible impédance d'entrée. Une haute impédance de sortie qui est fonction exclusivement de R_C .
- **Inconvénients** : son utilisation est souvent limitée à la haute fréquence, pas d'amplification en courant.

a) Caractéristiques

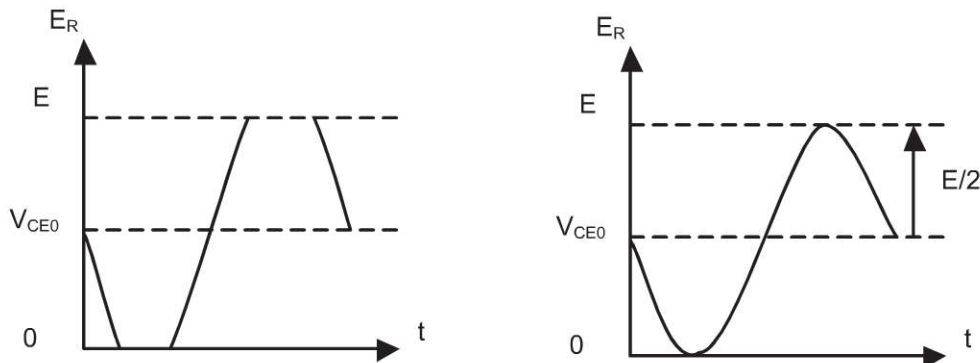
L'amplificateur de classe A consomme une partie de la puissance fournie par l'alimentation pour sa propre polarisation et une partie pour l'amplification du signal.

On suppose que le signal de sortie est sinusoïdal et donc non saturé, ce qui donne une tension de sortie constituée d'une composante continue $V_{C0} = E/2$ et d'une partie alternative « utile » d'amplitude $E_R/2$.

Donc $E_R(t) = E/2 + E_R/2 \times \sin(\omega t)$. Le courant traversant le transistor est la somme d'une partie continue (polarisation) valant $I_{C0} = E/2 \times R_{ER}$ et d'une partie alternative sinusoïdale

$$I_{C0}(t) = I_{CM}/2 \times \sin(\omega t) = E_R / (2 \times R_{ER}) \times \sin(\omega t)$$

► Amplitude maximale de E_R



Elle correspond à la limite de la tension d'alimentation de l'amplificateur, à savoir E . Dans le cas d'un dépassement nous obtenons un écrêtage de la tension E_R

► Puissance moyenne fournie par l'alimentation

$$Pf = \frac{1}{T} \int_0^T E \times (I_{C0} + I_{C0}(t)) dt = \frac{1}{T} \int_0^T E \times \left(\frac{E}{2 \times R_{ER}} + \frac{E_R}{2 \times R_{ER}} \sin(\omega t) \right) dt = \frac{E^2}{2R_{ER}}$$

La puissance fournie est indépendante de la puissance demandée par la charge.

► Puissance à dissiper par le transistor

$$\begin{aligned} Pd &= \frac{1}{T} \int_0^T V_{CE}(t) \times (I_{C0} + I_{C0}(t)) dt \\ &= \frac{1}{T} \int_0^T \left[\left(\frac{E}{2} - E_R \sin(\omega t) \right) \times \left(\frac{E}{2R_{ER}} + \frac{E_R}{R_{ER}} \sin(\omega t) \right) \right] dt = \frac{E^2}{4R_{ER}} - \frac{E_R^2}{2R_{ER}} \end{aligned}$$

La puissance à dissiper est maximale lorsque la tension E_R est nulle, donc aucun signal à amplifier. Elle est minimale lorsque la tension E_R est au bord de la saturation, c'est-à-dire pour une amplitude de $E/2$.

► Puissance utile pour l'amplification

$$\begin{aligned} Pu &= \frac{1}{T} \int_0^T E_R(t) \times (I_{C0} + I_{C0}(t)) dt \\ &= \frac{1}{T} \int_0^T \left[\left(\frac{E}{2} + E_R \sin(\omega t) \right) \times \left(\frac{E}{2R_{ER}} + \frac{E_R}{R_{ER}} \sin(\omega t) \right) \right] dt = \frac{E^2}{4R_{ER}} + \frac{E_R^2}{2R_{ER}} \end{aligned}$$

Deux termes apparaissent, le terme fonction de E correspond à la partie continue de la puissance absorbée. Cette partie n'est pas la partie utile du signal. On désire avant tout amplifier un signal variable. C'est alors le 2^e terme fonction de E_R qui est reflet de l'amplification.

► Rendement en régime variable

Comme nous l'avons dit précédemment, seule la partie variable du signal nous intéresse. On peut alors définir le rendement comme suit :

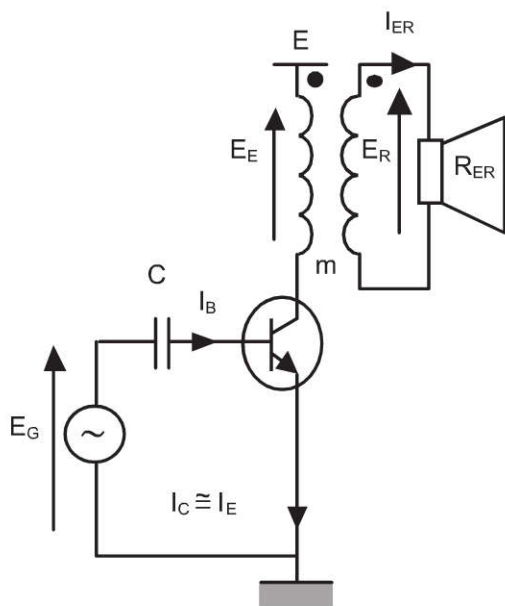
$$\eta = \frac{P_u}{P_f} = \frac{\frac{E_R^2}{2 \times R_{ER}}}{\frac{E^2}{2 \times R_{ER}}} = \frac{E_R^2}{E^2}$$

Nous avons vu que $E_R = E/2$ au maximum si l'on ne veut pas de distorsion du signal si le point de repos est bien au milieu de la droite de charge. Ce qui nous amène à un rendement théorique de 25 %. La caractéristique de charge dynamique, qui représente l'évolution du point de fonctionnement pour le régime variable, se confond avec la caractéristique de charge.

Remarques

- En réalité, le rendement est encore bien inférieur à 25 % puisque $E_R < E/2$, car la caractéristique I_C fonction de V_{CE} à I_B donné n'est pas linéaire (voir caractéristique en début de fiche). Il faut alors limiter l'excursion de E_R .
- Il est possible d'améliorer le rendement de ce type de montage en utilisant un transformateur (voir application).

3. EN PRATIQUE



On désire amplifier un signal E_G de façon à ce qu'il devienne audible. Pour cela on utilise un amplificateur de classe A et un haut parleur d'impédance Z_{ER} que l'on supposera purement résistif à la fréquence du signal à amplifier $Z_{ER} = R_{ER}$.

On définit m comme le rapport de transformation entre les valeurs efficaces des grandeurs alternatives comme :

$$m = \frac{E_R}{E_E} = \frac{I_C}{I_{ER}}$$

La résistance vue par le primaire vaut $R = R_{ER} / m^2$ puisque $E_R = R_{ER} \times I_{ER}$

$$E_R = m \times E_E \text{ et } I_{ER} = I_C / m$$

L'insertion du transformateur permet de diviser l'impédance de la charge par m^2 .

$$\begin{cases} V_{CE}(t) = V_{CE0} + V_{CEM} \times \sin(\omega t) \\ I_C(t) = I_{C0} + I_{CM} \times \sin(\omega t) \end{cases}$$

On désire obtenir un point de repos de $V_{CE0} = 8 \text{ V}$ pour $I_{C0} = 5 \text{ mA}$, avec $E = 16 \text{ V}$. L'impédance du haut parleur est de 8Ω . On désire obtenir la valeur de m permettant une amplitude maximale du signal de sortie sans distorsion.

L'équation de la droite de charge concernant le point de repos :

$$I_{C0} = \frac{E - V_{CE0}}{R}$$

D'où $R = (16 - 8) / 5 \cdot 10^3 = 1,6 \text{ k}\Omega$ ce qui donne :

$$m = \sqrt{\frac{R_{ER}}{R}} = 0,1$$

On peut remarquer, d'après la caractéristique ci-après, que l'amplitude $V_{CE}(t)$ sera inférieure à 8 V si l'on ne veut pas de distorsion du signal. Cette caractéristique n'étant pas à l'échelle, il faudrait s'attendre à une amplitude de 7 V max.

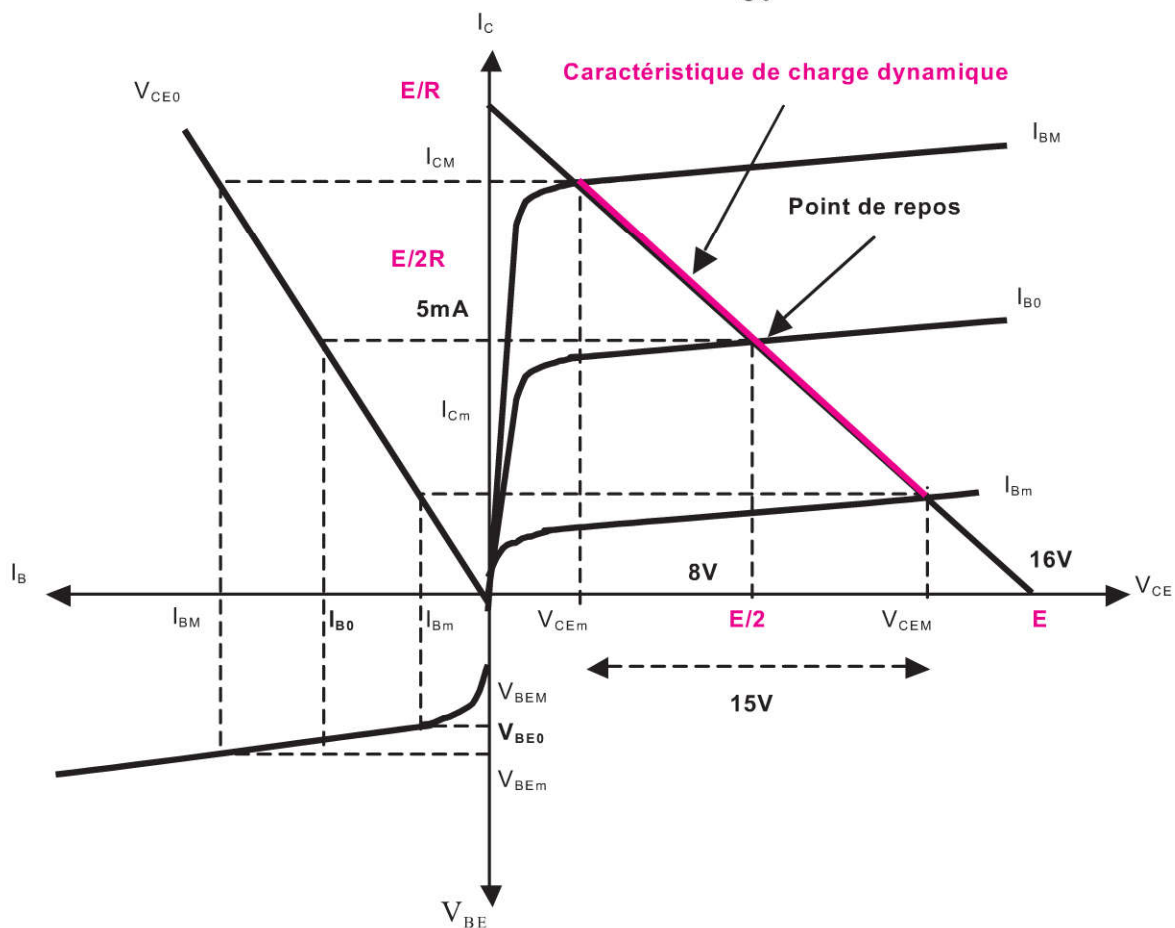
La puissance utile pour la charge vaut :

$$P_u = \left(\frac{V_{CEM}}{\sqrt{2}} \right)^2 \times \frac{1}{R}$$

La puissance fournie par l'alimentation vaut $P_f = E \times I_{C0} = 16 \times 5 \cdot 10^{-3} = 80 \text{ mW}$.

► Rendement

$$\eta = \frac{Pu}{Pf} = \frac{V_{CEM}^2}{2 \times E \times R \times I_{C0}}$$



a) Conclusion

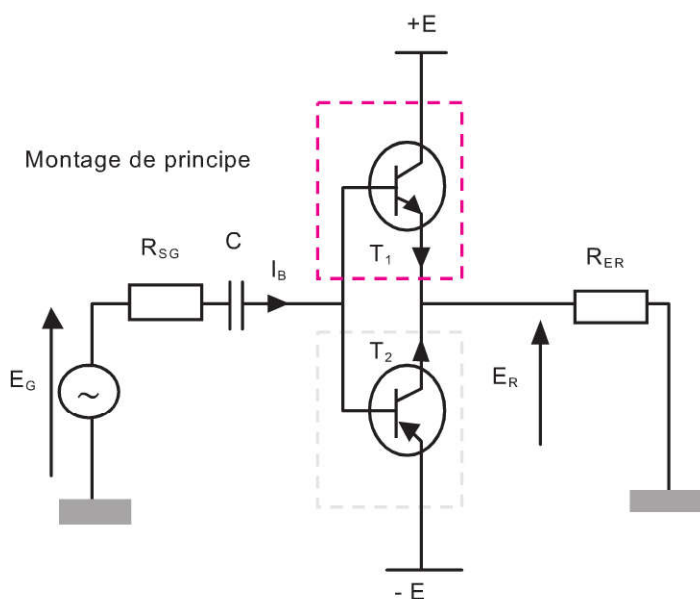
- **Avantages** : excellente linéarité car utilise pleinement la caractéristique de sortie du transistor. Seule une légère distorsion apparaît du fait de la pente naturelle de IC en fonction de VCE à IB donné. Simple de conception.
- **Inconvénients** : rendement faible. Taille du dissipateur très conséquent. Grande consommation d'énergie.
- **Application** : étages préamplificateurs, des amplificateurs audio, amplificateurs hautes fréquences à large bande ainsi que des oscillateurs hautes fréquences.

41 Amplificateur de classe B

Mots clés

Puissance, fournie, perdue, utile, conduction, alternée.

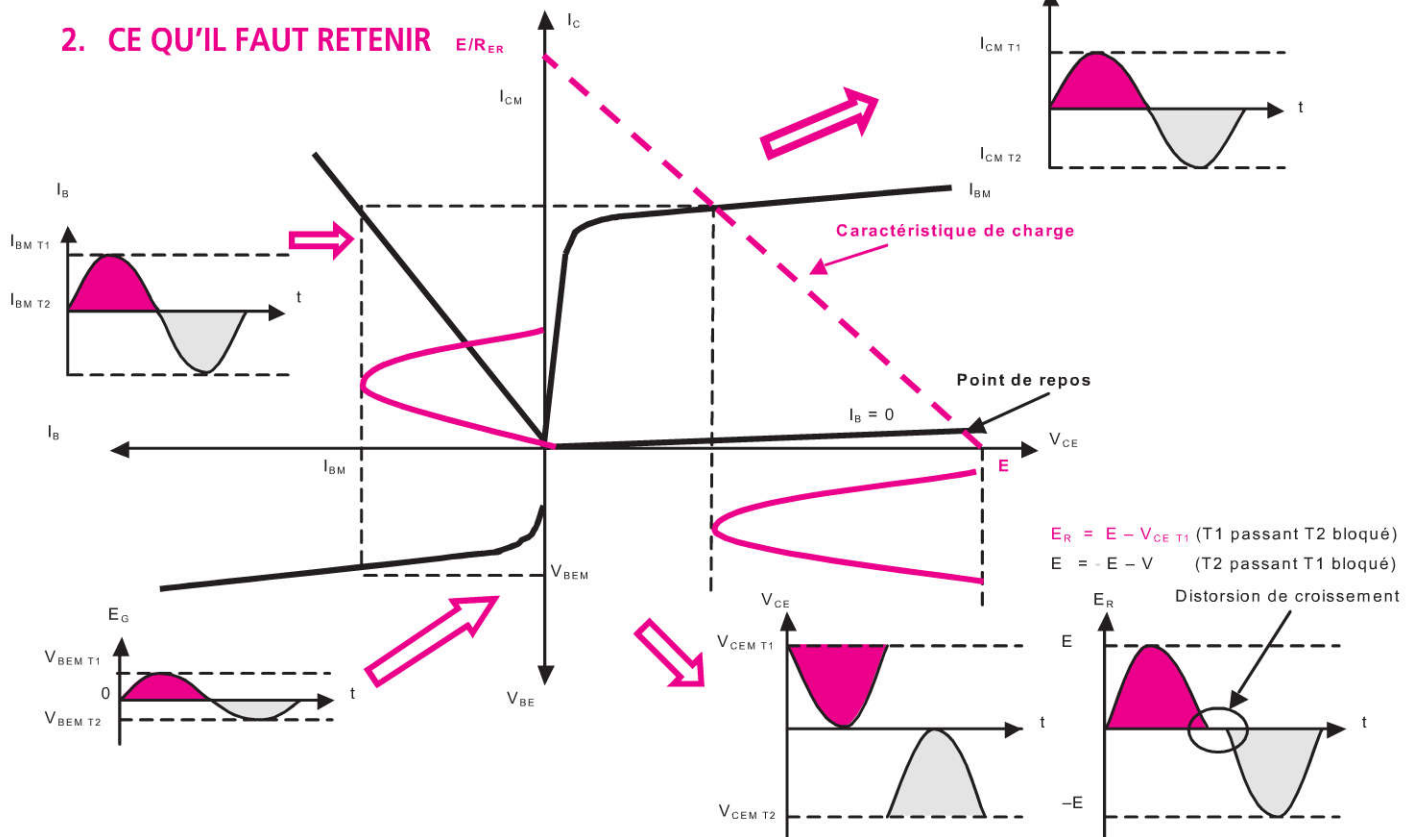
1. EN QUELQUES MOTS



L'inconvénient majeur du fonctionnement des amplificateurs de classe A est leur forte consommation en l'absence de signal à amplifier. Les classes B se comportent comme des amplificateurs de classe A à commande complémentaire.

Ceci permet de choisir un point de repos à courant nul lorsqu'aucun signal n'est appliqué. Ainsi, les transistors T1 (NPN) et T2 (PNP) fonctionnent de façon complémentaire durant une demi-période chacun. On suppose que les caractéristiques de T1 et T2 sont identiques.

2. CE QU'IL FAUT RETENIR

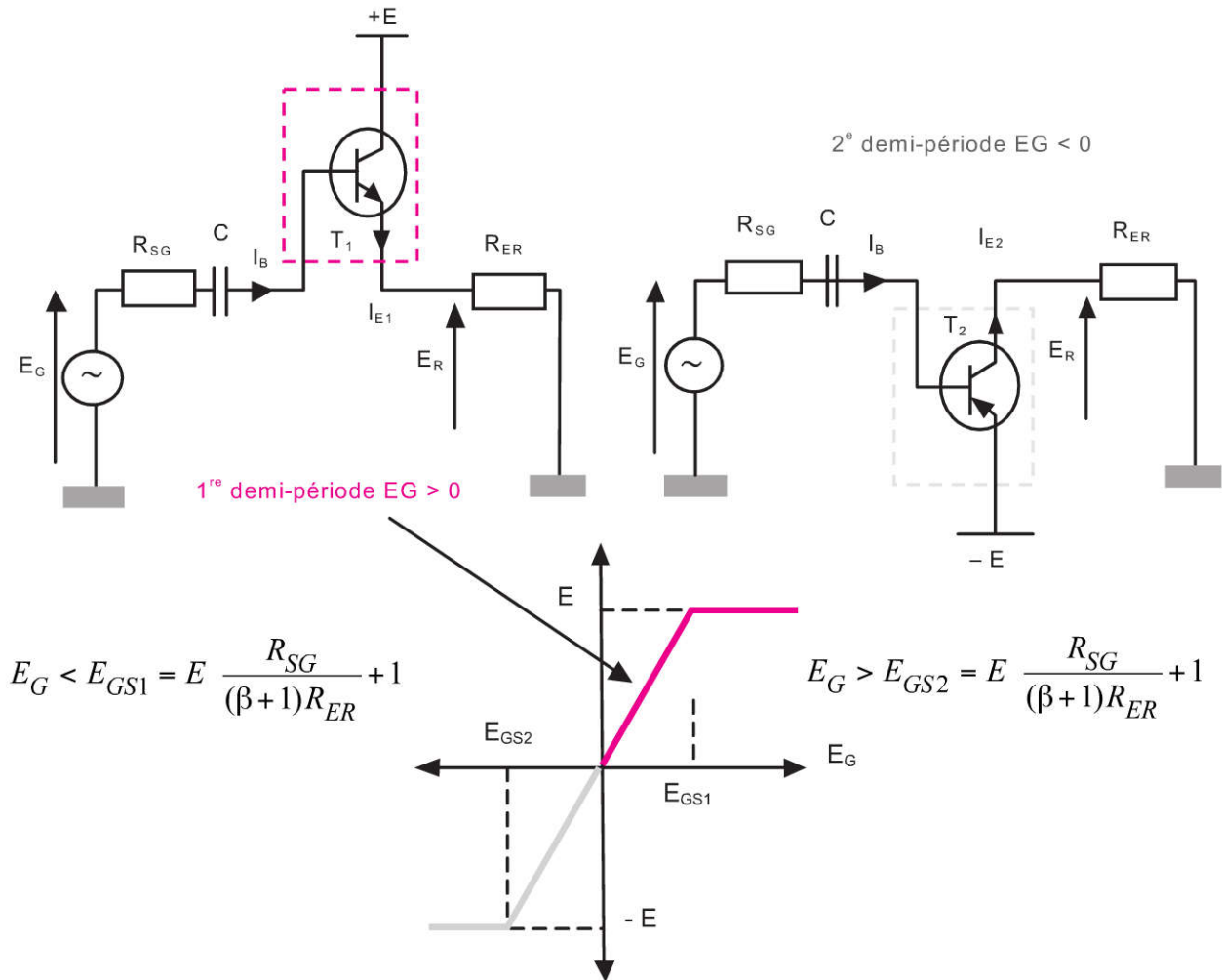


a) Caractéristiques

Le montage présenté est un « push-pull », les émetteurs étant reliés, $V_{BE\ T1} = V_{BE\ T2}$.

b) Seuil d'amplification

Le montage amplifie, sans écrêtage, chaque demi-période à la condition que la tension E_G reste inférieure ou supérieure à un seuil noté E_{GS} . On néglige l'effet de p sur le modèle des transistors, et nous sommes en régime linéaire, $I_E = (\beta + 1) \times I_B$. $\beta_1 = \beta_2 = \beta$



La tension en sortie vaut $E_R(t) = R_{ER} \times I_{E1} = \frac{R_{ER} \times (\beta + 1) \times E_{G \max} \sin(\omega t)}{R_G + (\beta + 1) \times R_{ER}}$

$$G_V = \frac{E_R}{E_G} = \frac{R_{ER} \times (\beta + 1)}{R_G + (\beta + 1) \times R_{ER}}$$

► **Puissance moyenne fournie par l'alimentation**

Chaque transistor fournit un courant I_E sous une tension E :

$$Pf(t) = E \times I_{E1}(t) - E \times I_{E2}(t) \text{ et } Pf = \frac{1}{T} \int_0^T pf(t) dt$$

$$Pf = \frac{1}{T} \left(\int_0^{T/2} E \times \frac{E_{R \max}}{R_{ER}} \sin(\omega t) dt - \int_{T/2}^T E \times \frac{E_{R \max}}{R_{ER}} \sin(\omega t) dt \right) = \frac{2 \times E \times E_{R \max}}{\pi \times R_{ER}}$$

► **Puissance utile**

Si le signal E_R est sinusoïdal d'amplitude maximale E alors la puissance utile vaut :

$$Pu = \left(\frac{E_R}{\sqrt{2}} \right)^2 \times \frac{1}{R_{ER}} \quad Pu_{\max} = \left(\frac{E}{\sqrt{2}} \right)^2 \times \frac{1}{R_{ER}}$$

► **Puissance moyenne dissipée par transistor**

$$Pd_{total} = Pf - Pu = \frac{Pd_{total}}{2} = \frac{E \times E_R}{\pi \times R_{ER}} - \frac{E_R^2}{4 \times R_{ER}}$$

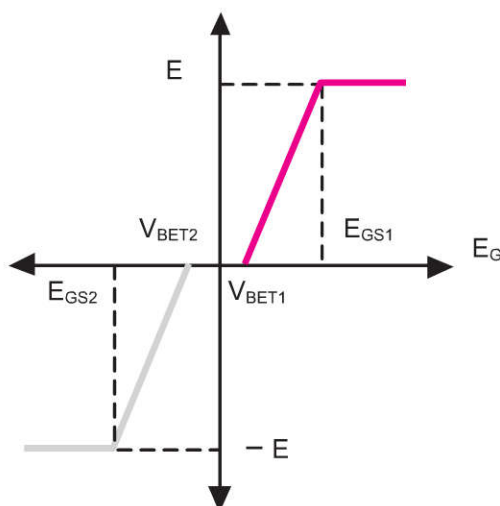
La puissance maximale à dissiper dans un transistor est fonction de E_R . Elle n'est pas max pour $E_R = E$ mais pour $E_R = 2 \times E/\pi$. Il faut dériver l'expression de $Pd/2$ en fonction de E_R et faire l'égalité avec 0. La puissance max dissipée par un transistor vaut alors :

$$Pd_{T1 \max} = Pd_{T2 \max} = E^2 / (\pi^2 \times R_{ER}).$$

► **Rendement**

$$\eta = \frac{Pu}{Pf} = \frac{\pi \times E_R}{4 \times E} \times 100 \text{ le rendement est max lorsque } E_R = E, \eta = 78,5 \%$$

► **Distorsion de croisement du signal de sortie**



Le défaut principal de ce type de montage est lié à la conduction de T_1 et T_2 lorsque la tension E_G est proche de 0. Les jonctions émetteur-base ne sont passantes que si la tension E_G est supérieure à la tension de seuil.

En réalité, la caractéristique E_R en fonction de E_G présente un palier autour de 0, tant que $E_G < \approx 0,7 \text{ V}$. Les transistors ne peuvent pas conduire, d'où une tension de sortie E_R nulle dans cette zone.

On retrouve alors l'allure de la tension E_R qui est donnée en première page de cette fiche qui correspond à une distorsion dite de bas niveau (c'est lorsque E_G est proche de 0,7 V qu'il y a problème).

► **Amélioration par utilisation de diodes (figure 1)**

Le problème étant lié à la tension de seuil des transistors, l'idée serait d'annuler cette tension en y ajoutant un composant en parallèle dont la tension de seuil serait identique. C'est cette fonction que remplissent les diodes D_1 et D_2 . Les résistances R_1 et R_2 limitent le courant dans le sens direct des diodes et délivrent le courant I_B de chaque transistor. Il est à noter que le réglage pratique de ces deux résistances n'est pas toujours évident à effectuer.

La boucle de tension surlignée en rouge nous donne $E_G - R_{SG} \times I_G + U_{D1} - V_{BE1} - E_R = 0$

$U_{D1} \cong V_{BE1}$ et $R_{SG} \times I_G \ll E_R$ et E_G . D'où $E_R \cong E_G$.

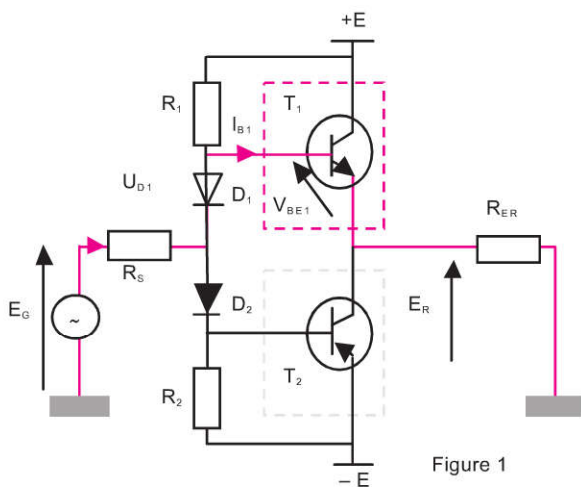


Figure 1

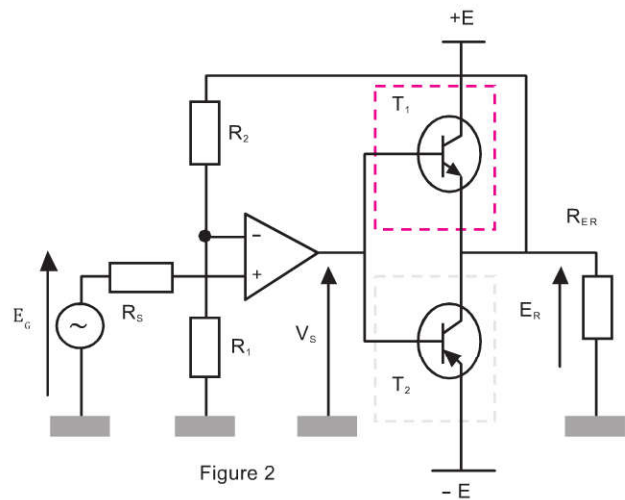


Figure 2

► **Amélioration par utilisation d'un ALI (figure2)**

Un ALI monté en amplificateur non inverseur permet d'obtenir un gain en tension lorsqu'un transistor est passant réglable par les valeurs de R_1 et R_2 tel que : $G_V = \frac{E_R}{E_G} = \frac{R_1 + R_2}{R_1}$

Pour que l'étage de puissance fonctionne, il faut que la tension V_S soit supérieure à la tension de seuil de la jonction base - émetteur $\cong 0,7$ V.

Grâce à la très grande amplification de l'ALI ($Ad \cong 10^5$), $V_S = Ad (V^+ - V^-)$. Ce qui donne $E_{Gmin} = V_{BE}/Ad$ soit 7 μ V. Cela permet de supprimer la distorsion à bas niveau.

c) Conclusion

- **Avantages** : rendement beaucoup élevé que pour un amplificateur de classe A. Taille du radiateur beaucoup moins importante.
- **Inconvénient** : discontinuité du signal de sortie (distorsion de croisement). Cette distorsion peut être atténuée au prix de l'utilisation de plusieurs composants.
- **Applications** : étage de sortie des amplificateurs continu utilisés dans : les boucles d'asservissement linéaire, les générateurs de fonction, les amplificateurs en circuits intégrés, ainsi que dans l'étage de sortie de la majorité des amplificateurs audio.

42 Transistors à effet de champ

Mots clés

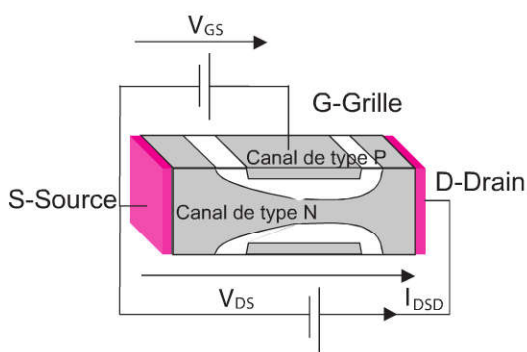
Canal, pincement, seuil, zone ohmique, zone active, signe V_{GS} .

1. EN QUELQUES MOTS

Les transistors à effet de champ (TEC) regroupent deux grandes familles : les **JFET** (*Junction Field Effect Transistor*) et le **MOSFET** (*Metal Oxyde Semiconductor Field Effect Transistor*).

2. CE QU'IL FAUT RETENIR

a) JFET

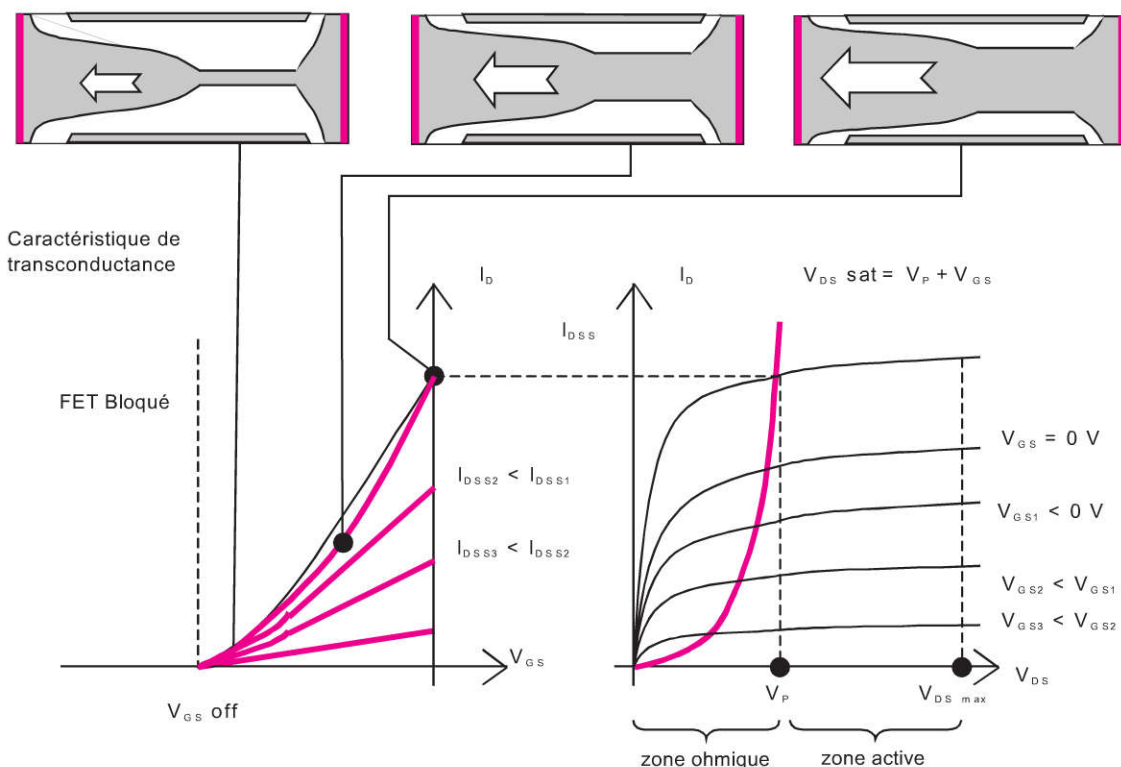


Si on alimente sous une tension V_{GS} , on vient agrandir ou rétrécir le canal permettant la circulation des électrons entre le drain et la source. On parle de canal N lorsque le dopage de ce canal est de type N, et de canal P dans le cas inverse. Selon le type de canal, le signe de V_{GS} permet de modifier sa taille.

C'est donc la tension V_{GS} qui va permettre le réglage du courant circulant entre le drain et la source.

Visualisation du courant de drain I_D en fonction V_{GS} pour un canal de type N. Plus la tension V_{GS} diminue, plus le canal se rétrécit et s'oppose au passage du courant de Drain.

On note $V_{GS\ off}$ la tension grille de fermeture, V_P la tension de pincement = $-V_{GS\ off}$.



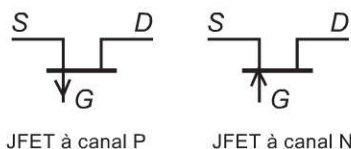
Ce que l'on peut dire du fonctionnement

- La tension V_{GS} doit passer d'une valeur nulle à une tension négative (pour un canal N) pour faire débiter le transistor de presque rien à son courant maximal. Il faut donc une assez grande excursion du signal d'entrée pour procurer une grande variation de I_D .
- Le FET se commande en tension, l'impédance d'entrée est très élevée, aucun courant (ou presque) n'est consommé sur l'étage précédent (normal car Z_e est très grand).
- Cette caractéristique de Drain ressemble beaucoup à la caractéristique de collecteur du transistor bipolaire. On y trouve une zone ohmique et une zone active.
- Dès que l'on a quitté la zone ohmique, la région active est très plate et le transistor est utilisable sur une grande plage de V_{DS} . Au-delà d'une certaine valeur de V_{DS} , la jonction entre dans la zone de claquage et le courant s'envole.
- Lorsque le transistor est « normalement passant », I_D est maximal pour $V_{GS} = 0$, et diminue lorsqu'on augmente V_{GS} (en valeur absolue). I_D est nul lorsque V_{GS} dépasse une valeur limite $V_{GS\ off}$.

Canal P : $V_{GS} > 0 \Leftrightarrow$ la charge positive sur la grille repousse les trous.

Canal N : $V_{GS} < 0 \Leftrightarrow$ la charge négative sur la grille repousse les électrons.

Les symboles et équations de fonctionnement sont les suivants :



Le trait associé à la grille peut être décalé, comme ici, ou centré sur le composant.

Équation de caractéristique de transconductance : $V_{DS} > V_{DS\ sat} : I_D = I_{DSS} \left(1 - \frac{V_{GS}}{V_T} \right)^2$

Équation de la caractéristique de sortie, en zone ohmique :

$$I_D = 2 \frac{I_{DSS}}{V_T^2} \left(V_{GS} - V_T - \frac{V_{DS}}{2} \right) \times V_{DS}$$

pour $V_{GS} < V_{GS\ off}$, $I_D \approx 0$, transistor bloqué.

pour $V_{GS} > 0$, le courant I_G augmente rapidement (zone non utilisée).

Les caractéristiques constructeur (valeurs numériques d'un FET de type 2N5484) correspondent aux paramètres suivants :

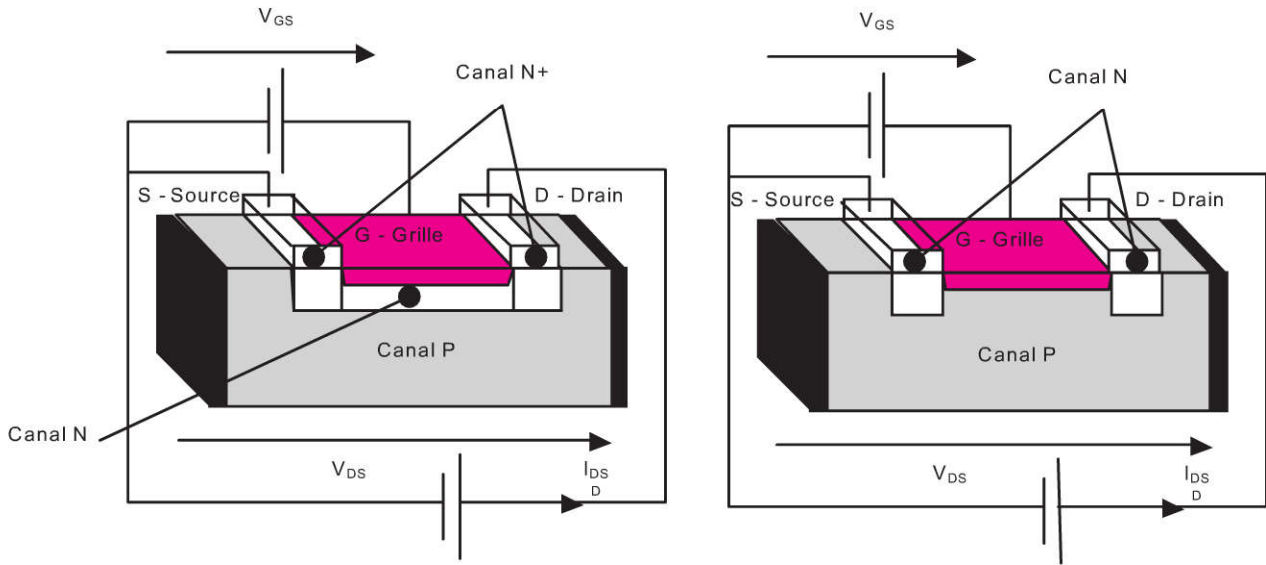
- V_{GS} : tension entre grille et source. C'est la tension de polarisation du FET qui va commander le courant qui va s'écouler du drain vers la source (-25 V).
- V_{DS} : tension entre drain et source, identique au concept V_{CE} d'un transistor bipolaire.
- $V_{GS\ off}$: tension pour laquelle le transistor est bloqué et ne débite plus ($-0,3\text{ min à }-3\text{ V max}$).
- I_D : courant qui circule entre Drain et source.
- I_{DSS} : courant de drain quand la grille est court-circuitée, c'est à dire $V_{GS} = 0$. Ce courant représente le courant maximum que peut débiter le transistor ($1\text{ min à }5\text{ mA max}$).

b) MOSFET

Il existe deux types de MOS : MOS à appauvrissement (à canal N) et à enrichissement (à canal P), les MOS à enrichissement de type P n'étant pas fabriqués.

Les deux figures ci-dessous nous placent dans le cas de canaux de type N. La figure 1 concerne les MOS appauvris et la figure 2 les enrichis. On parle toujours de canal N pour les

MOS appauvris même si le canal N n'existe pas physiquement, contrairement aux enrichis. Il se crée naturellement lors de l'application de la tension V_{GS} mais de façon moins aisée, ce qui se traduit par la nécessité d'appliquer une tension $V_{GS} > 0$.



Ce que l'on peut dire du fonctionnement

Les symboles et équations de fonctionnement sont les suivants :

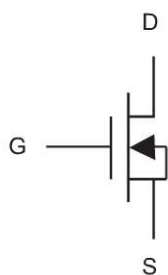
- Équation de caractéristique de transconductance $V_{DS} > V_{DS\text{ sat}}$:

$$I_D = I_{DSS} \left(1 - \frac{V_{GS}}{V_T} \right)^2$$

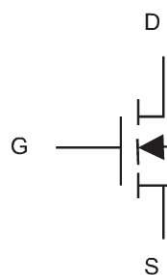
- Équation de la caractéristique de sortie, en zone ohmique :

$$I_D = 2 \frac{I_{DSS}}{V_T^2} \left(V_{GS} - V_T - \frac{V_{DS}}{2} \right) \times V_{DS}$$

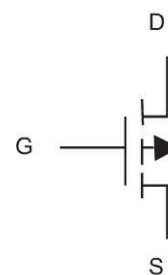
pour $V_{GS} < V_S$, $I_D \approx 0$, transistor bloqué, $V_{GS} - V_S =$ « tension d'attaque de grille ».



Canal N appauvri
 $I_D > 0$



Canal N enrichi
 $I_D > 0$



Canal P enrichi
 $I_D < 0$

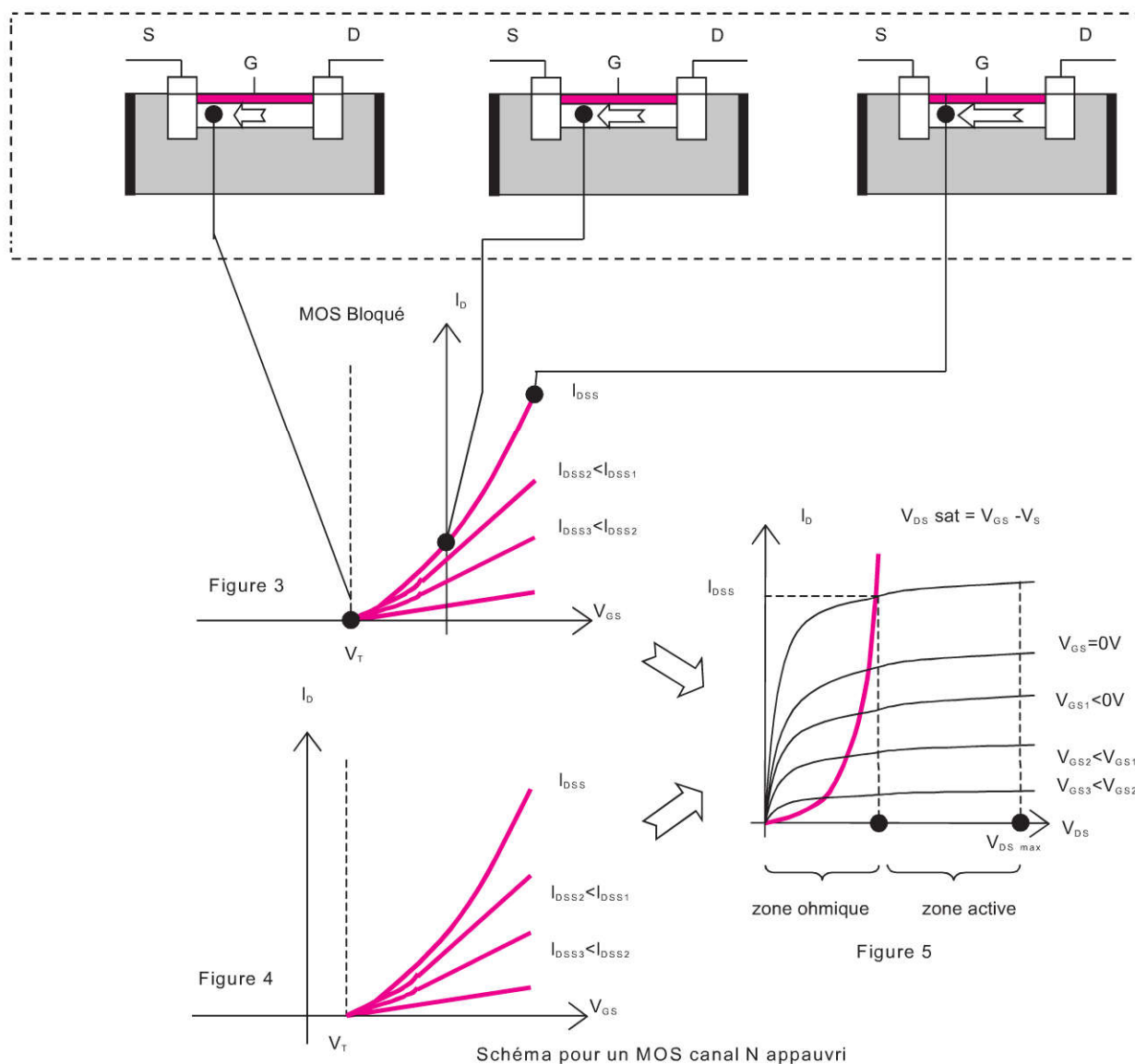
- Pour un MOS de type canal N à appauvrissement il est nécessaire d'appliquer une tension $V_{GS} < V_T$ (négative), pour obtenir un courant $I_D = 0$. En effet si aucune tension V_{GS} n'est appliquée, le canal N existe par construction et permet donc la circulation des électrons de la Source vers le Drain. Ceci nous donne la caractéristique de la figure 3.
- À l'inverse, les MOS de type canal N à enrichissement ne peuvent permettre la circulation des électrons de la Source vers le Drain, tant qu'une tension $V_{GS} > V_T$ (positive) n'est pas

appliquée. Le canal doit alors se construire en repoussant les trous du canal P. On appelle V_T la tension de seuil (threshold). Nous obtenons alors les caractéristiques de la figure 4.

- Qu'il soit à enrichissement ou à appauvrissement, la caractéristique I_D en fonction de V_{DS} d'un TEC est identique. L'allure de cette caractéristique (figure 5) est aussi la même que celle des JFET.
- Pour un MOS à enrichissement, $I_D = 0$ lorsque $V_{GS} = 0$, et augmente dès que V_{GS} dépasse une valeur seuil V_T .

S'il est de type canal P : $V_s < 0 \Leftrightarrow$ la charge négative sur la grille attire les trous.

S'il est de type canal N : $V_s > 0 \Leftrightarrow$ la charge positive sur la grille attire les électrons.



43 Polarisation des TEC

Mots clés

Caractéristique entrée, sortie, courant de gâchette nul, dispersion.

1. EN QUELQUES MOTS

Dans cette fiche, nous abordons quelques montages de polarisation de transistors à effet de champ. En raison de l'absence de courant de grille, des circuits de polarisation nouveaux (par rapport au cas bipolaire) peuvent être utilisés. Il est à noter que la polarisation doit permettre une évolution de I_D et V_{DS} dans la zone ohmique, ce qui n'est pas forcément représenté sur les caractéristiques qui suivent.

2. CE QU'IL FAUT RETENIR

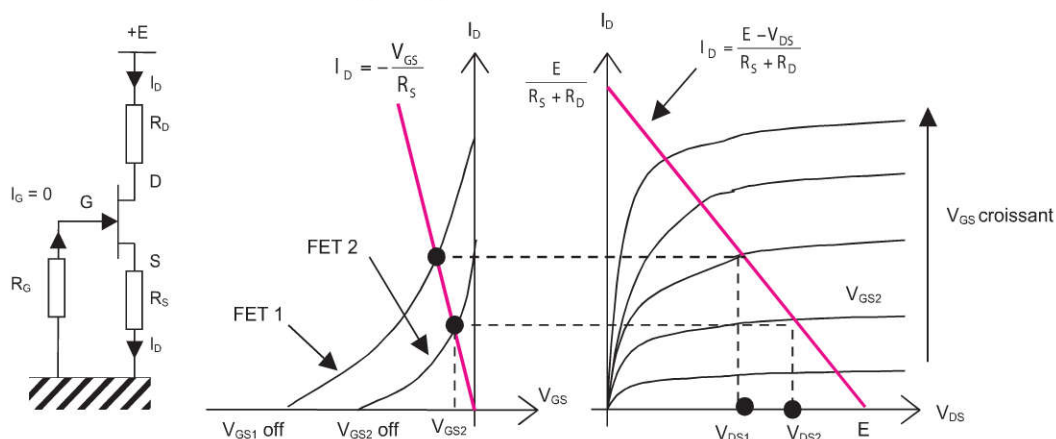
a) Auto-polarisation (exemple d'un JFET canal N)

Le circuit d'auto-polarisation permet de fixer automatiquement la tension V_{GS} par l'intermédiaire de la résistance R_S . La résistance R_G est nécessaire pour relier la grille à la masse. Sa valeur va déterminer l'impédance d'entrée des amplificateurs. On sait que le courant I_G est quasi nul.

Les transistors à effet de champ souffrent d'une dispersion importante (FET 1 remplacé par le FET 2 d'une même famille) de fabrication. Il en résulte une incertitude sur les paramètres I_{DSS} et V_{GSoff} . I_D dépend de V_{GS} , donc c'est la tension présente à la source du transistor qui va le polariser : $V_{GS} = V_G - V_S = 0 - I_D \times R_S$ donc $V_{GS} = -I_D \times R_S$.

La caractéristique d'entrée est une droite $I_D = -\frac{V_{GS}}{R_S}$

Caractéristique de sortie $I_D = \frac{E - V_{DS}}{R_S + R_D}$

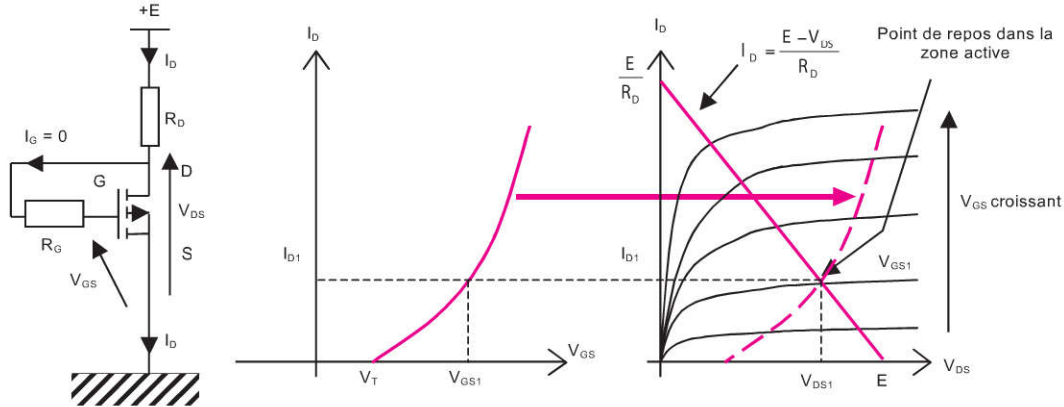


- **Inconvénient** : pour les mêmes valeurs de R_G , R_D et R_S , la dispersion entraîne des variations notables de V_{DS} et I_D . Points de repos différents.
- **Avantage** : permet de stabiliser la valeur de I_D . Si I_D augmente, $R_S \times I_D$ augmente, la tension V_{GS} augmente et devient plus négative, d'où diminution de I_D .

b) Polarisation par réaction de drain (exemple d'un MOS canal N enrichi)

La caractéristique d'entrée et la caractéristique I_D en fonction de V_{GS} sont les mêmes puisque $I_G \equiv 0$.

Caractéristique de sortie :
$$I_D = \frac{E - V_{DS}}{R_S + R_D}$$



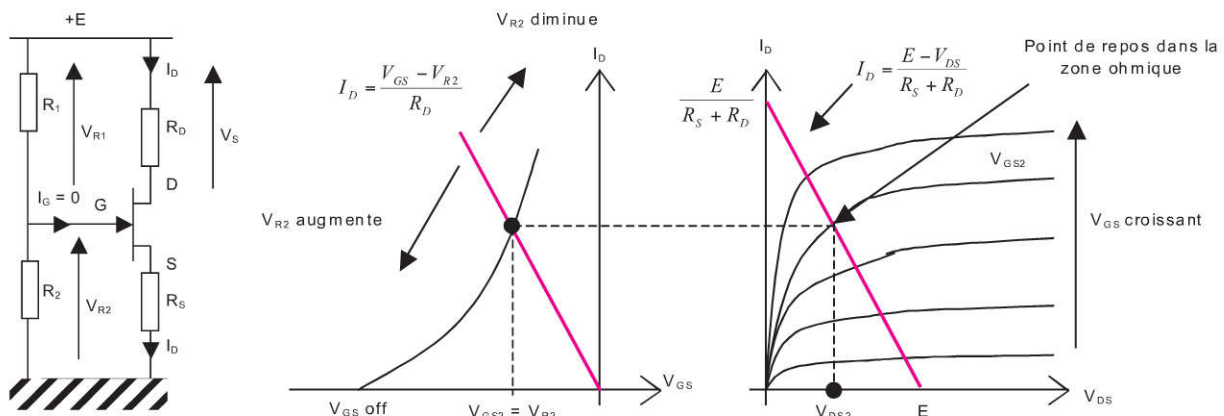
- **Inconvénient** : la dispersion entraîne des variations notables de V_{DS} et I_D . Seule la caractéristique du transistor impose le point de repos. Pas de possibilité de réajuster le point de repos avec les résistances extérieures.
- **Avantage** : n'utilise que deux résistances (ordre de grandeur de R_G : plusieurs centaines k Ω).

c) Polarisation en pont (exemple d'un JFET canal N)

Dans ce mode de polarisation, c'est la tension V_{R2} qui va fixer la tension V_{GS} . Puisque $I_G \equiv 0$ V_{R2} est déduite d'un simple diviseur de tension.

$$V_G = V_{R2} = \frac{R_2}{R_1 + R_2} \times E \quad V_S = R_D \times I_D \quad V_{GS} = V_{R2} - R_S \times I_D$$

Caractéristique d'entrée : $I_D = \frac{E - V_{DS}}{R_D + R_S}$ Caractéristique de sortie : $I_D = \frac{V_{GS} - V_{R2}}{R_S}$



3. EN PRATIQUE

On désire le point de repos suivant : $V_{GS} = -1$ V, $I_D = 5$ mA, $V_{DS} = 3$ V, $E = 10$ V. Quelles sont les valeurs de R_1 , R_2 , R_D et R_S ? Pour cela, il faut s'imposer certaines valeurs, effectuer les calculs puis vérifier si le point est à peu près au milieu de la zone ohmique.

On s'impose $R_1 = R_2 = 20$ k Ω d'où $V_{R2} = 5$ V. D'où $R_D = (V_{GS} - V_{R2}) / I_D = -(-1 - 5) / 5 \cdot 10^{-3} = 1\,200$ Ω . $R_S = (E - V_{DS} - R_D \times I_D) / I_D = (10 - 3 - 1\,200 \times 5 \cdot 10^{-3}) / 5 \cdot 10^{-3} = 200$ Ω .

44 Modèle du transistor à effet de champ

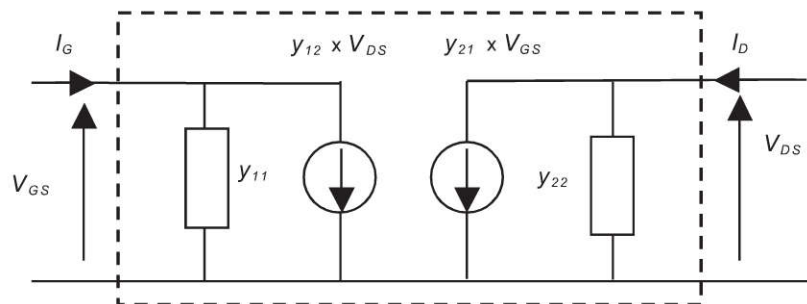
Mots clés

Paramètres, modèle, impédance, facteur d'amplification, conductance, capacités, modèles.

1. EN QUELQUES MOTS

Le transistor TEC peut être vu, comme le transistor bipolaire, comme un quadripôle Q dont l'entrée serait V_{GS} et les sorties V_{DS} et I_D . Il est important de signaler que les grandeurs $I_G = 0$, I_D , V_{GS} , V_{DS} sont **ici des grandeurs variables sinusoïdales**.

2. CE QU'IL FAUT RETENIR



Le modèle proposé s'applique pour des **signaux sinusoïdaux évoluant autour d'un point de repos** fixé par un montage de polarisation. L'amplitude de ces signaux est faible devant les grandeurs de repos. Ceci permet de rester dans les zones linéaires de fonctionnement des différentes caractéristiques et c'est ce qui explique aussi les valeurs élevées ou très faibles des paramètres.

a) Les paramètres du transistor TEC

Représentation des paramètres admittances
$$\begin{cases} I_G = y_{11} \cdot V_{GS} + y_{12} \cdot V_{DS} \\ I_D = y_{21} \cdot V_{GS} + y_{22} \cdot V_{DS} \end{cases}$$

Puisque I_G est quasiment nul du fait même de sa constitution, les paramètres y_{11} et y_{12} sont nuls.

Détermination graphique des paramètres
$$\Delta I_D = y_{21} \cdot \Delta V_{GS} + y_{22} \cdot \Delta V_{DS}$$

On note Δ pour signifier qu'il s'agit d'une variation autour du point de repos fixé par le montage polarisant. On suppose les variations d'amplitudes suffisamment faibles pour considérer les caractéristiques I_D fonction de V_{GS} et I_D fonction de V_{DS} linéaires. Pour avoir un ordre d'idée des grandeurs, nous utiliserons les caractéristiques du transistor JFET à canal N 2N5668.

b) Le paramètre y_{21} du transistor TEC (transconductance ou pente)

Si la tension V_{DS} est constante, lorsque la tension d'entrée varie d'une quantité ΔV_{GS} , le courant de drain varie d'une quantité $\Delta I_D = y_{21} \times \Delta V_{GS}$. On nomme y_{21} la pente ou la transconductance du transistor.

$$y_{21} = \left(\frac{\Delta I_D}{\Delta V_{GS}} \right)_{\Delta V_{DS}=0} = \left(\frac{\Delta I_D}{\Delta V_{GS}} \right)_{V_{DS}=cst}$$

nom constructeur $y_{21} = g_{fs}$ (*forward transfert conductance*).

ordre de grandeur min-max = de 3 à 7,5 mS

c) Le paramètre y_{22} du transistor TEC (résistance dynamique de sortie).

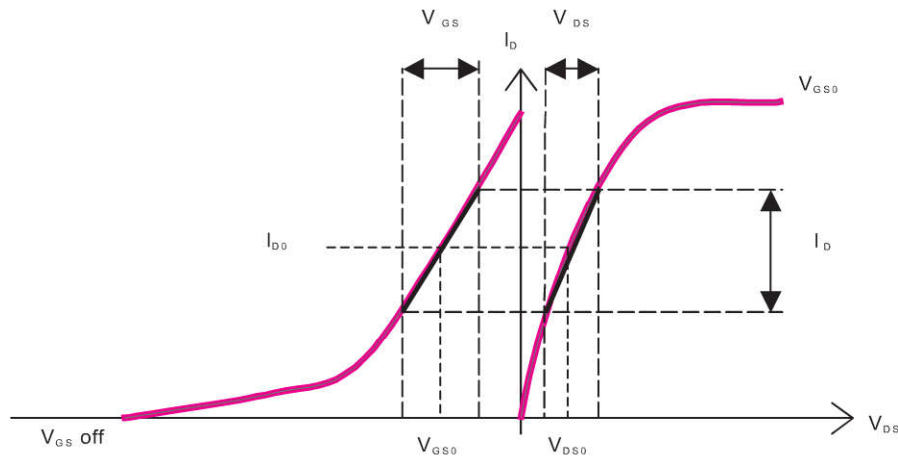
Si la tension V_{GS} est constante, lorsque la tension de sortie varie d'une quantité ΔV_{DS} , le courant de drain varie d'une quantité $\Delta I_D = y_{22} \times \Delta V_{DS}$. On nomme y_{22} l'inverse de la résistance dynamique de sortie.

$$y_{22} = \left(\frac{\Delta I_D}{\Delta V_{DS}} \right)_{\Delta V_{GS}=0} = \left(\frac{\Delta I_D}{\Delta V_{DS}} \right)_{V_{GS}=cst}$$

nom constructeur $y_{22} = g_{os}$ (output conductance) ou $R_{DS(on)} = 1/y_{22}$.

ordre de grandeur ($R_{DS(on)}$) min-max = de 100 à 500 Ω

Il est à noter que les échelles ne sont pas respectées et que ces caractéristiques ne font qu'illustrer les paramètres exprimés précédemment.



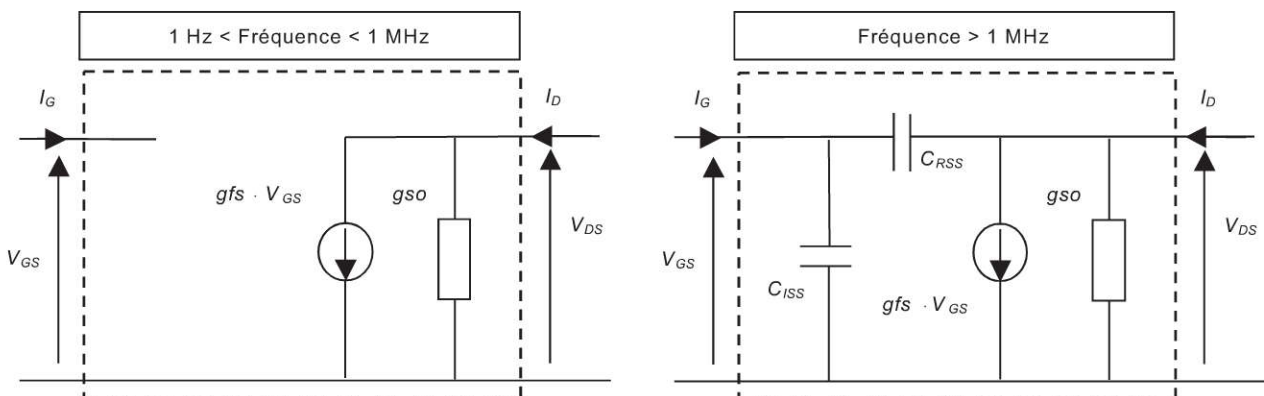
Remarque

Comme le transistor bipolaire, il y a dérive des paramètres en fonction de la température et aussi en fonction du courant de drain I_D . Ainsi, la valeur « g_{os} » diminue si I_D augmente, et « g_{fs} » augmente si I_D augmente, tout ceci à température donnée.

3. EN PRATIQUE

a) Modèle du transistor TEC « basse » et « haute » fréquence

La limite de fréquence est difficile à définir dans le sens où l'on glisse d'un modèle « basse fréquence » vers un modèle « haute fréquence » de façon progressive.



Le modèle « haute » fréquence fait intervenir des capacités supplémentaires qui apparaissent avec l'augmentation de F . Ainsi, nous avons C_{ISS} (Input Capacitance) 3 à 4 pF pour notre JFET et C_{RSS} (Reverse Transfer Capacitance) de 0,6 à 0,9 pF.

45 Montage à source commune

Mots clés

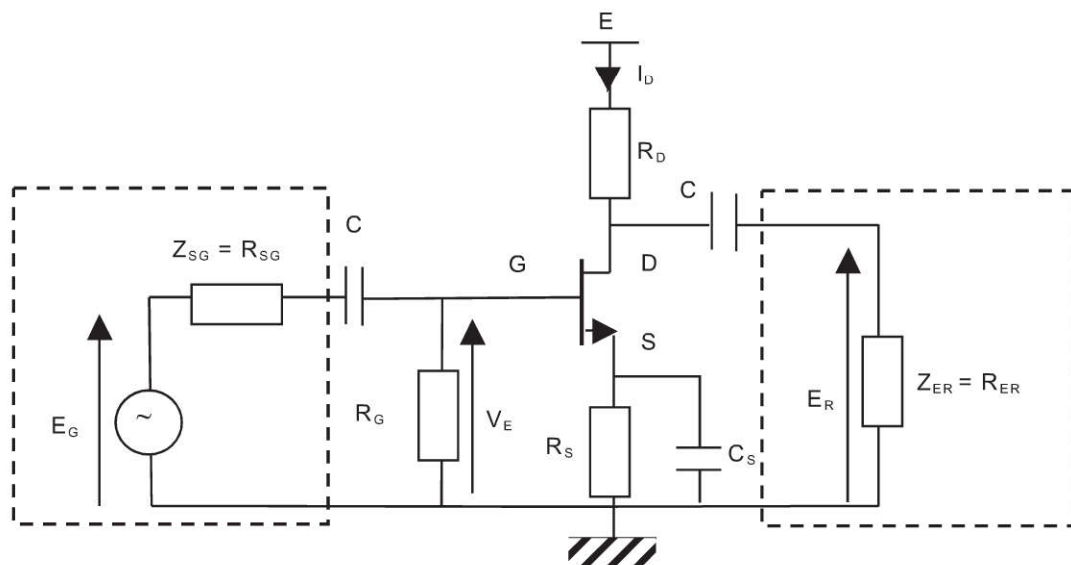
Impédance d'entrée grande, faible gain inverse, intermédiaire.

1. EN QUELQUES MOTS

Le montage à amplificateur à effet de champ source commune est typiquement utilisé en tant qu'amplificateur de tension. Sa résistance d'entrée est faible, ce qui implique qu'il n'a pas besoin d'être associé avec un montage adaptateur d'impédance.

2. CE QU'IL FAUT RETENIR

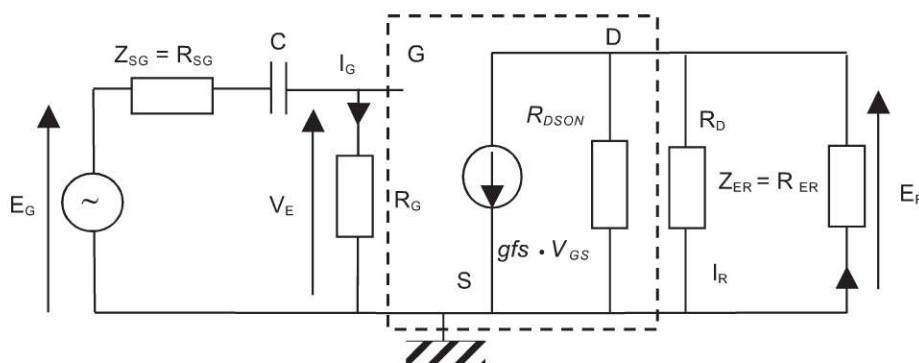
On rappelle que les grandeurs en tensions et courants sont sinusoïdales. Les capacités C sont des condensateurs de découplage permettant de séparer grandeurs du point de repos et grandeurs alternatives. Les valeurs de C dépendent des impédances vues par le signal alternatif et peuvent être toutes différentes, C_S étant une capacité permettant de découpler la résistance de source.



a) Schéma équivalent pour les petits signaux

La résistance R_G a une valeur importante, de l'ordre du $M\Omega$. C'est pour cette raison qu'il n'y a pas de courant dans la grille. C'est donc R_G qui conditionne l'impédance d'entrée du montage.

En appliquant le modèle du TEC à notre JFET à canal N pour les « basses » fréquences, nous obtenons :



Équations déduites du schéma : $V_E = E_G \times \frac{R_G}{R_{SG} + R_G} \cong E_G$

$$R_T = \frac{R_{DS(on)} \times R_D}{R_{DS(on)} + R_D}$$

$$E_R = -I_{ER} \times R_{ER}$$

b) Performances de ce montage

► **Gain en tension G_V :** $G_V = \frac{E_R}{V_E} = -g_{fs} \times \frac{R_T \times R_{ER}}{R_T + R_{ER}}$

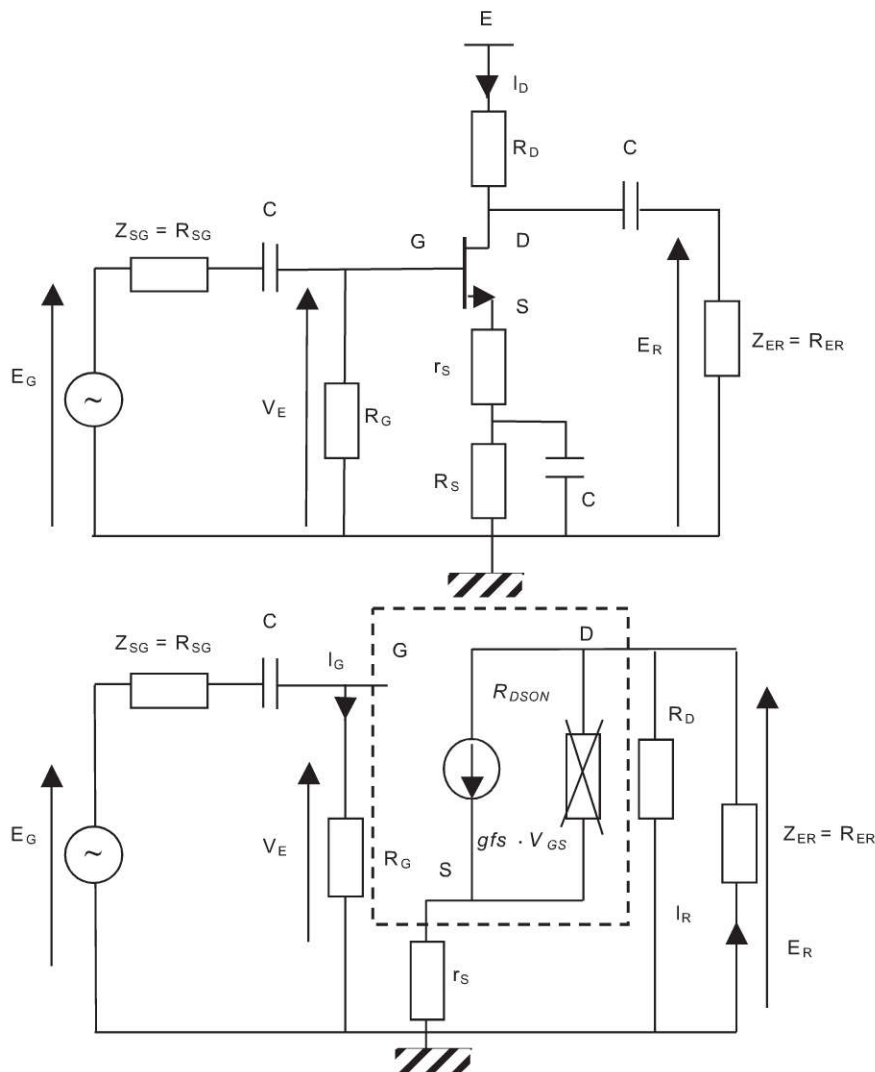
Il y a opposition de phase entre la tension de sortie et la tension d'entrée.

► **Impédance d'entrée Z_{EA} :** $Z_{EA} = \frac{V_E}{I_G} = R_G$

► **Impédance de sortie Z_{SA} :** $Z_{SA} = \frac{E_R}{I_R} = R_T$

L'impédance de sortie est déduite lorsque le générateur de tension E_G d'entrée est un court-circuit et que la charge est déconnectée.

c) Variante au montage à source commune



Afin d'obtenir une performance du gain en tension indépendante des caractéristiques du transistor, on peut insérer une résistance r_S en série avec la Source.

La résistance r_S introduit une contre réaction $V_{GS} = V_G - r_S \times I_D$.

Le schéma équivalent pour les petits signaux fait apparaître la résistance r_S entre la Source et la masse. On peut négliger l'influence de $R_{DS\text{ON}}$, en la considérant comme infinie, si les caractéristiques du transistor le permettent.

Exemple

2N5484 a pour valeur $g_{os\text{ max.}} = 50 \mu S$ soit une résistance minimale de $20 \text{ k}\Omega$.

Dans ce cas, à charge déconnectée : ($I_R = 0$)

$$V_E = V_{GS} + r_S \times g_{fs} \times V_{GS}$$

$$E_R = -g_{fs} \times V_{GS} \times R_D$$

soit un nouveau gain en tension :

$$G_V = \frac{E_R}{V_E} = \frac{-R_D}{\frac{1}{g_{fs}} + r_S}$$

Lorsque r_S est bien supérieure à $1/g_{fs}$, le gain ne dépend plus des caractéristiques du transistor.

3. EN PRATIQUE

Depuis le schéma suivant, on donne les valeurs ci-dessous :

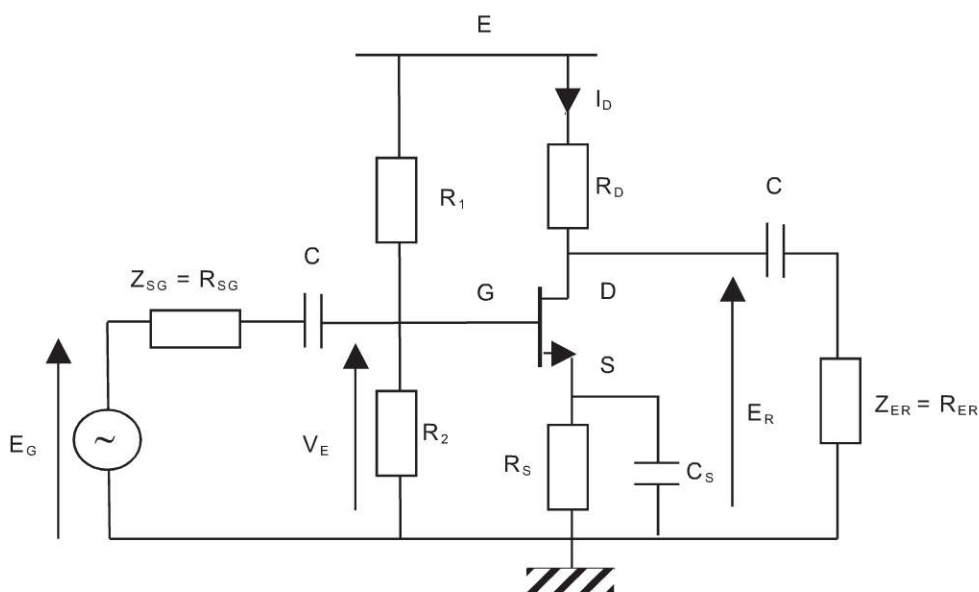
Amplitude de la tension crête à crête de $E_G = 10 \text{ mV}$, $R_G = 500 \Omega$, $R_1 = R_2 = 1 \text{ M}\Omega$, $E = 20 \text{ V}$, $R_D = 3,6 \text{ k}\Omega$, $R_S = 10 \text{ k}\Omega = R_{ER}$, un JFET 2N5484 donne un $g_{fs\text{ min.}}$ de 6 mS et $g_{s0} = 75 \mu S$ d'où $R_{DS\text{ON}} = 1 / 75 \cdot 10^{-6} = 13,3 \text{ k}\Omega$.

Le gain en tension vaut : $G_V = -g_{fs} \times (R_D // R_{DS\text{ON}} // R_{ER}) = -6 \cdot 10^{-3} \times 1,52 = 9,12$.

L'amplitude la tension E_R de sortie crête à crête vaut 91 mV en opposition de phase par rapport à E_G .

L'impédance d'entrée : $Z_{EA} = R_1 // R_2 = 500 \text{ k}\Omega$ qui $\gg$ devant R_G .

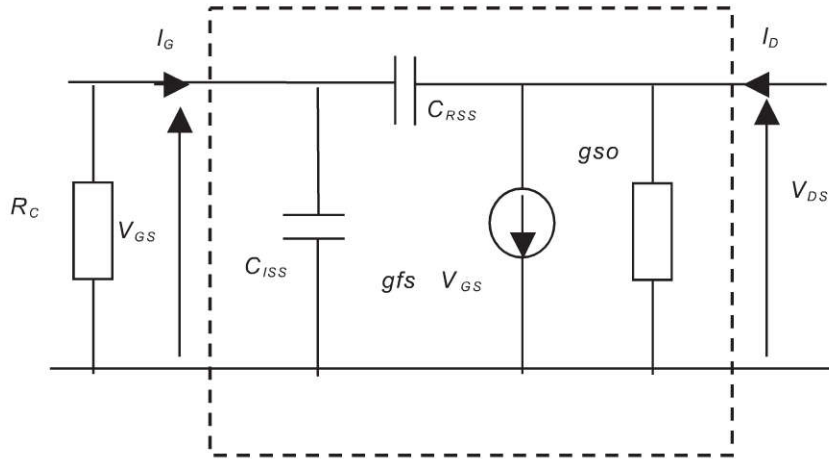
L'impédance de sortie $Z_{SA} = R_T = R_D // R_{DS\text{ON}} = 2,83 \text{ k}\Omega$.



Remarque

Pour des signaux dont la fréquence est supérieure au MHz, les capacités C_{SS} et C_{RSS} modifient considérablement les caractéristiques de ce montage. Ainsi, la capacité C_{SS} fait chuter considérablement l'impédance d'entrée : c'est l'avantage principal de ce montage.

Le courant I_G ne peut plus être considéré comme nul.

**a) En conclusion**

- **Avantage** : ce type de montage permet d'avoir une grande résistance d'entrée. Cette dernière est définie par les valeurs imposées à R_1 et R_2 , dans le cas d'une polarisation en pont, et par R_G en auto-polarisation.
- **Inconvénient** : le gain est fonction de g_{fs} qui reste faible (toujours inférieur à 10 mS).
- **Application** : ce montage est utilisé comme montage d'entrée intermédiaire.

46 Montage à drain commun

Mots clés

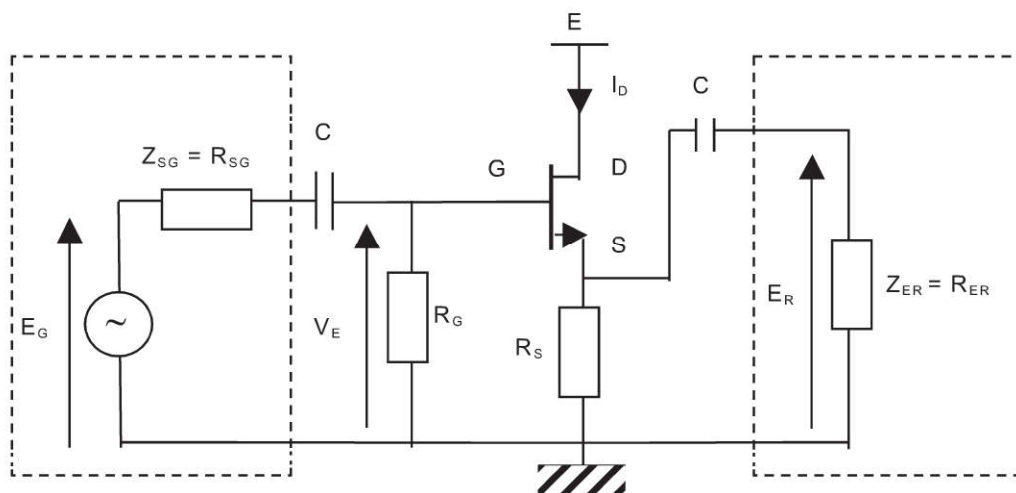
Gain < 1 , impédance de sortie faible, impédance d'entrée élevée, adaptation d'impédance.

1. EN QUELQUES MOTS

Le montage à transistor à effet de champ drain commun fait partie de la famille des systèmes adaptateurs d'impédance : leur gain est inférieur à l'unité (atténuation) et leur résistance d'entrée est faible.

a) Ce qu'il faut retenir

Les conditions sur les signaux sont identiques au montage à source commune.

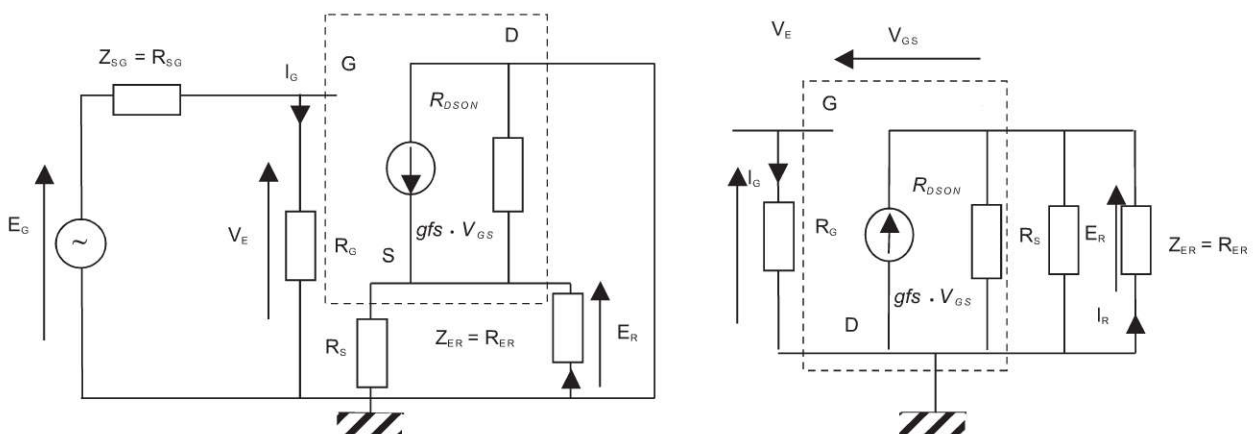


Caractéristique de la droite d'attaque : $V_{GS} = -R_S \times I_D$

Caractéristique de la droite de charge : $V_{DS} = E - R_D \times I_D$

b) Schéma équivalent pour les petits signaux

Les conditions sont identiques au montage à source commune.



Le schéma équivalent de droite est le schéma « brut », où l'on a simplement inséré le modèle du transistor au sein du montage. Le schéma équivalent de gauche est le même, dans une version plus retravaillée.

On s'aperçoit que le Drain et la Source sont inversés, et que les résistances R_{DSON} , R_S et R_{ER} sont en parallèle. Pour la tension V_{GS} , elle se trouve aux bornes de R_G et du générateur de courant. Équations déduites du schéma :

$$V_{GS} = V_E - E_R \quad R_T = \frac{R_{DSON} \times R_S}{R_{DSON} + R_S} \quad E_R = g_{fs} \times V_{GS} \times R_T$$

c) Performances de ce montage

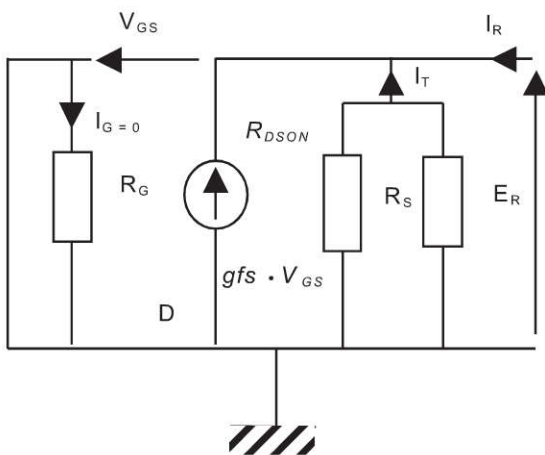
► **Gain en tension G_V** : $G_V = \frac{E_R}{V_E} = g_{fs} \times \frac{R_T}{g_{fs} \times R_T + 1} < 1$

Puisque $E_R = g_{fs} \times R_T \times (V_E - E_R)$ d'où $E_R (1 + R_T \times g_{fs}) = V_E \times R_T \times g_{fs}$

Contrairement au montage transistor bipolaire à collecteur commun, le gain à vide n'est pas très voisin de 1 car $R_T \times g_{fs}$ n'est pas beaucoup plus grand que 1.

► **Impédance d'entrée** : $Z_{EA} = \frac{V_E}{I_G} = R_G$

► **Impédance de sortie** : $\frac{E_R}{I_R} = \frac{1/g_{fs} \times R_T}{1/g_{fs} + R_T} = 1/g_{fs} // R_T$



L'impédance de sortie est déduite lorsque le générateur de tension E_G d'entrée est un court-circuit et que la charge est déconnectée.

Une tension E_R est appliquée et fait circuler un courant I_R . $I_R = -I_T - g_{fs} \times V_{GS}$ avec $V_{GS} = -E_R$ d'où :

$$I_R = \frac{E_R}{R_T} - g_{fs} \times V_{GS} = \frac{E_R}{R_T} + g_{fs} \times E_R$$

d) En pratique

Pour avoir une idée des ordres de grandeurs, nous imposons les valeurs suivantes aux résistances de notre montage. On suppose que ces valeurs permettent une polarisation correcte du JFET.

$R_G = 1 \text{ M}\Omega$, $R_S = R_{ER} = 2 \text{ k}\Omega$, $E = 10 \text{ V}$, amplitude crête à crête de $E_G = 100 \text{ mV}$, $R_{SG} = 500 \Omega$. $R_{DSON} = 10 \text{ k}\Omega$ et $g_{fs} = 5 \text{ mS}$.

$$R_T = 1,6 \text{ k}\Omega \quad G_V = 0,88$$

$$Z_{EA} = 1 \text{ M}\Omega \quad Z_{SA} = 177 \Omega$$

La tension de sortie E_R après amplification vaut 88 mV crête à crête.

e) Conclusion

- **avantages** : ce montage permet d'obtenir une impédance d'entrée importante et qui est directement imposable par le choix de R_G . La résistance de sortie est faible par rapport à un montage à source commune.
- **inconvénient** : le gain en amplification est < 1 . Ce montage est destiné à être utilisé comme montage d'entrée (adaptation d'impédance).

47 Amplificateur différentiel

Mots clés

Tension différentielle, mode commun, mode différentiel, transistor bipolaire, TRMC, modèle équivalent petits signaux.

1. EN QUELQUES MOTS

L'amplificateur différentiel est un circuit amplificateur de tension adapté en tension capable d'amplifier des tensions continues et alternatives. Idéalement, le signal de sortie est égal à la différence des signaux d'entrée multipliée par le gain.

2. CE QU'IL FAUT RETENIR

Nous allons commencer par appréhender les différentes grandeurs caractéristiques d'un amplificateur différentiel (AD). Nous verrons par la suite comment est constitué ce type d'amplificateur mais, pour l'instant, nous considérons l'AD comme une boîte noire constituée de deux entrées et de deux sorties. Chacune de ces accès étant référencé à la masse, nous avons donc un octopôle.

La tension d'entrée différentielle est définie comme suit : $v_{ed} = v_{e1} - v_{e2}$

La tension d'entrée en mode commun est $v_{mc} = \frac{v_{e1} + v_{e2}}{2}$

De ces deux équations, nous pouvons déduire les relations suivantes :

$$v_{e1} = \frac{v_{ed} + 2 \cdot v_{mc}}{2} \text{ et } v_{e2} = \frac{-v_{ed} + 2 \cdot v_{mc}}{2}$$

Les tensions de sorties sont fonction des entrées et de l'amplification du montage :

$$v_{s1} = A_1 \cdot v_{e1} + A_2 \cdot v_{e2} = \frac{A_1 - A_2}{2} \cdot v_{ed} + (A_1 + A_2) \cdot v_{mc}$$

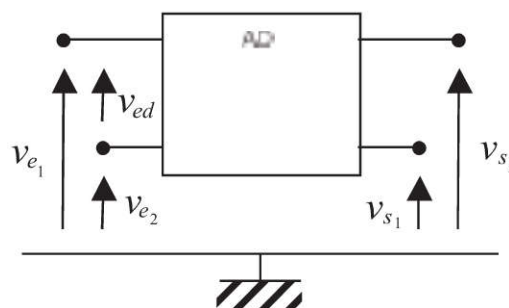
On pose $v_s = A_d \cdot v_{ed} + A_c \cdot v_{mc}$. A_d est le gain de l'amplificateur en mode différentiel et A_c le gain en mode commun.

Le taux de réjection en mode commun, ou TRMC, est un indicateur de la qualité d'un AD :

$$TRMC = \frac{A_D}{A_C}$$

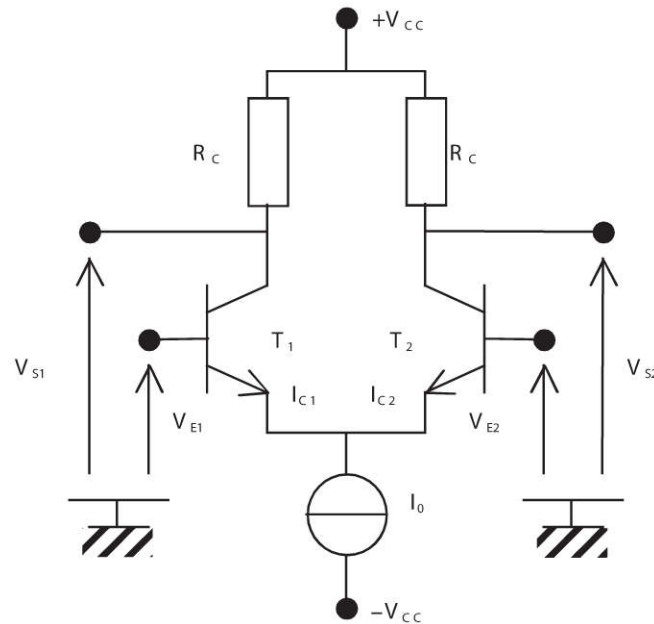
Le TRMC est généralement exprimé en décibels : $TRMC_{dB} = 20 \cdot \log \left(\frac{A_D}{A_C} \right)$

Pour un AD idéal, on considère que : $A_2 = A_1$. Par conséquent, le gain en mode commun est nul et le gain en mode différentiel est égal à A_1 , le TRMC tend vers l'infini.



3. EN PRATIQUE

L'amplificateur différentiel est basé sur deux transistors bipolaires T_1 et T_2 identiques (fiche 33). Le montage est polarisé par une alimentation symétrique $+V_{CC}$, $-V_{CC}$.



Les deux branches de l'amplificateur étant identiques, la loi des nœuds implique :

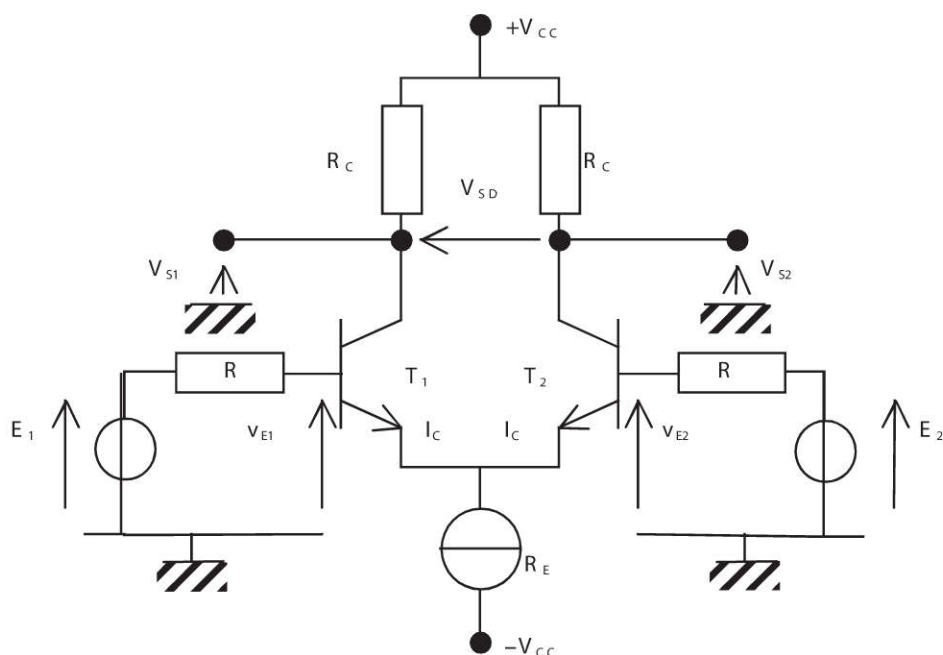
$$I_{C1} + I_{C2} = I_0 = cte$$

Par conséquent, en appliquant la loi des mailles entre $+V_{CC}$ et les entrées, on trouve :

$$V_{S1} = V_{S2} = V_{CC} - R_C \cdot I_0 / 2$$

a) Étude statique

Le montage est excité par deux générateurs de Thévenin de résistances internes identiques R et de tension E_1 et E_2 .



Chaque transistor bipolaire obéit à la loi exponentielle (fiche 34) :

$$I_{C1} = I_S \cdot e^{\frac{V_{BE1}}{U_T}} \text{ et } I_{C2} = I_S \cdot e^{\frac{V_{BE2}}{U_T}} \text{ avec } U_T \approx 27 \text{ mV à température ambiante (300 K)}$$

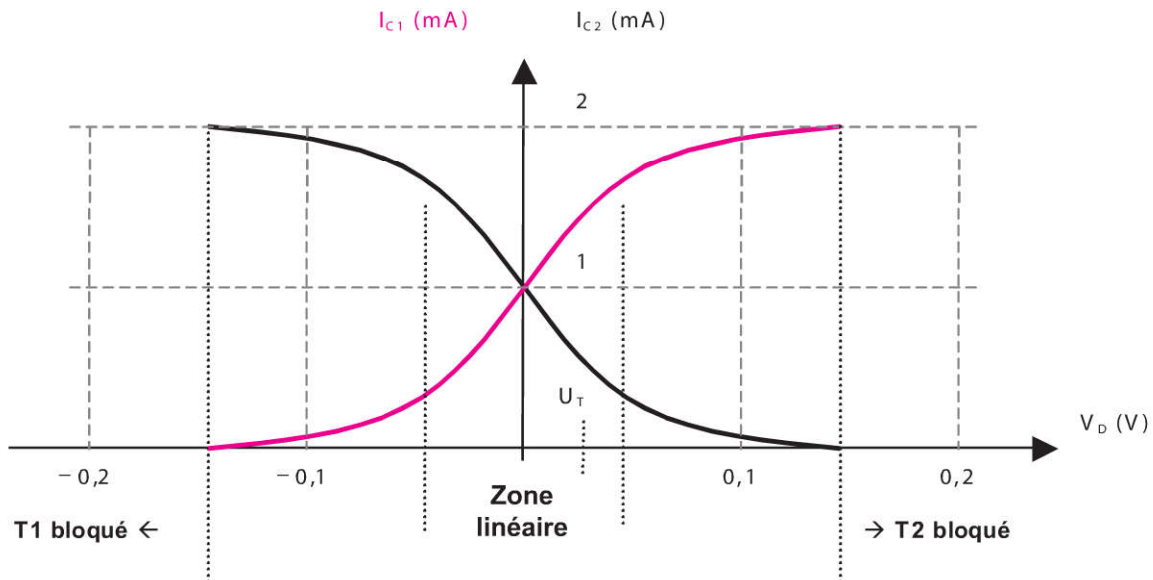
Soit la tension d'entrée différentielle V_D telle que :

$$V_D = V_{E1} - V_{E2} = V_{BE1} - V_{BE2}$$

$$\frac{I_{C1}}{I_{C2}} = e^{\frac{V_{BE1} - V_{BE2}}{U_T}} = e^{\frac{V_D}{U_T}} \approx 1 + \frac{V_D}{U_T} \text{ pour } V_D \ll U_T \text{ (27 mV)}$$

Avec les relations précédentes, on obtient les équations liant les courants de collecteur et la tension différentielle :

$$I_{C1} = \frac{I_0}{1 + e^{\frac{V_D}{U_T}}} \text{ et } I_{C2} = \frac{I_0}{1 + e^{-\frac{V_D}{U_T}}}$$



En sortie du montage, la tension différentielle V_{SD} s'écrit :

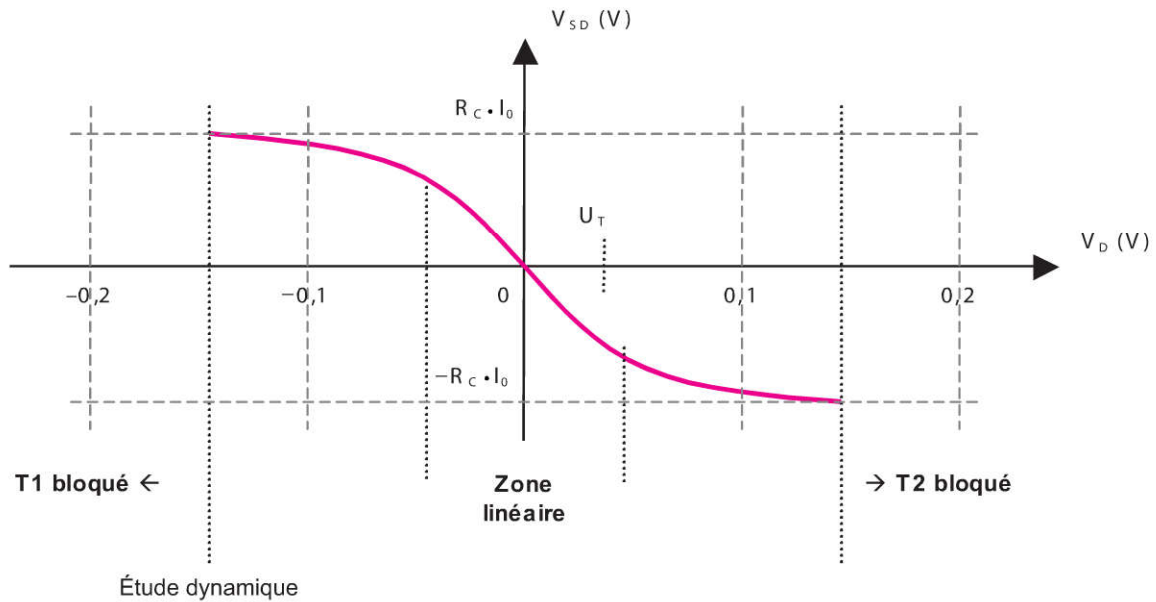
$$V_{SD} = R_C \cdot I_{C2} - R_C \cdot I_{C1} = R_C \cdot (I_{C2} - I_{C1})$$

En remplaçant I_{C1} et I_{C2} par leur expression, on trouve :

$$V_{SD} = R_C \cdot I_0 \cdot \frac{1 - 1 - \frac{V_D}{U_T}}{1 + 1 + \frac{V_D}{U_T}} = -R_C \cdot I_0 \cdot \frac{V_D}{2 \cdot U_T + V_D} \approx -R_C \cdot I_0 \cdot \frac{V_D}{2 \cdot U_T} \text{ avec } V_D \ll U_T$$

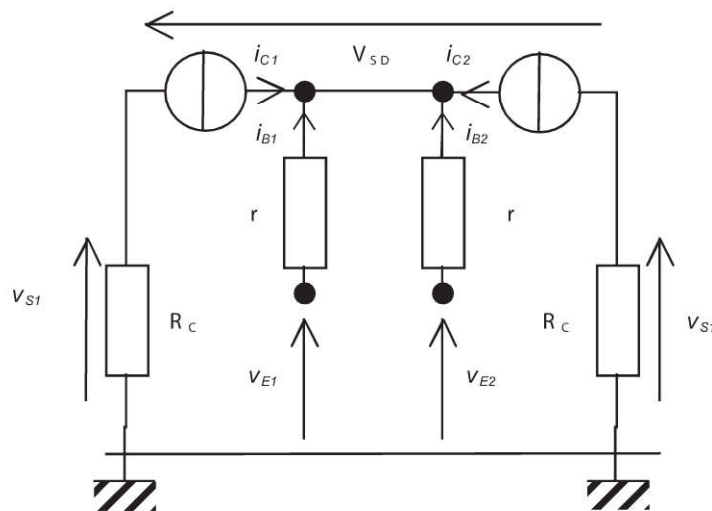
Nous pouvons en déduire l'expression du gain différentiel A_D dans la zone de fonctionnement linéaire :

$$A_D \approx -\frac{R_C \cdot I_0}{2 \cdot U_T}$$



b) Étude dynamique

L'amplificateur fonctionne en petits signaux (tensions d'entrée alternatives de faible amplitude). Le schéma équivalent du montage devient :



Les tensions différentielles d'entrée et de sortie en petit signaux sont déduites en appliquant la loi des mailles sur le schéma petits signaux :

$$v_D = v_{E1} - v_{E2} = r \cdot (i_{B1} - i_{B2}) \text{ et } v_{SD} = v_{S1} - v_{S2} = \beta \cdot R_C \cdot (i_{B1} - i_{B2})$$

Nous pouvons alors exprimer la tension différentielle de sortie en fonction de l'entrée et obtenir ainsi la fonction de transfert du montage :

$$v_{SD} = -\frac{\beta \cdot R_C}{r} \cdot v_D$$

Sachant que $i_{B1} = -i_{B2}$ et que $v_D = r \cdot (i_{B1} - i_{B2}) = 2 \cdot r \cdot i_{B1}$, on obtient :

$$v_{S1} = -v_{S2} = -\frac{\beta \cdot R_C}{2 \cdot r} \cdot v_D$$

Dans le cas idéal où le montage est parfaitement symétrique (composants « égaux »), le TRMC est infini. En réalité, la dispersion des valeurs de composants, notamment du β des transistors, conduit à avoir du gain en mode commun.

48 Amplificateur opérationnel

Mots clés

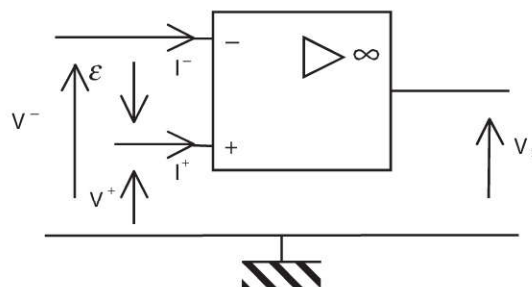
AOP idéal, AOP réel, courant de polarisation, gain en boucle ouverte, max slew rate.

1. EN QUELQUES MOTS

L'amplificateur opérationnel (AOP) est un circuit amplificateur à deux entrées et une sortie. Il est basé sur le principe des amplificateurs différentiels.

2. CE QU'IL FAUT RETENIR

L'AOP idéal est un modèle dont nous nous servirons dans la [fiche 50](#) « Montages à amplificateurs opérationnels ».



Le gain en mode différentiel (ou gain en boucle ouverte) est infini :

$$A_d = \frac{V_s}{\varepsilon} = \infty \Rightarrow \varepsilon = 0 \text{ V}$$

La résistance d'entrée est infinie entre :

- l'entrée + et la masse
- l'entrée - et la masse
- l'entrée + et l'entrée -

Ces hypothèses impliquent que I^+ et I^- sont nuls.

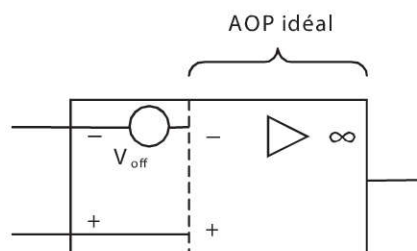
La résistance de sortie est nulle.

La bande passante de l'AOP est infinie, c'est-à-dire qu'un AOP idéal peut amplifier les tensions continues jusqu'à une fréquence infinie.

3. EN PRATIQUE

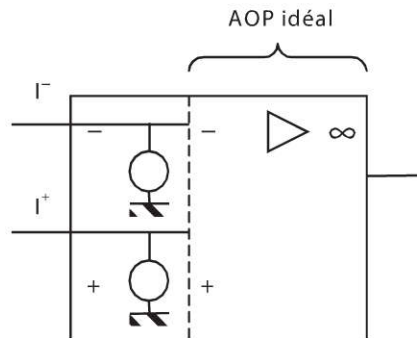
En pratique, l'amplificateur opération idéal n'existe pas. Nous parlons alors d'AOP réel.

On considère pour celui-ci la présence d'une tension de décalage ou d'« offset » V_{off} . Par conséquent, $\varepsilon \neq 0 \text{ V}$ implique que A_d est une valeur finie.

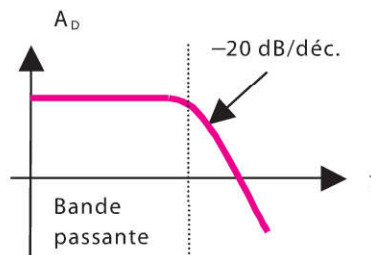


On considère également que les courants de polarisation I^+ et I^- ne sont pas nuls. À partir de ces courants, on définit :

- le courant de polarisation moyen $I_{\text{pol}} = \frac{I^+ + I^-}{2}$
- le courant de décalage (ou d'offset) $I_{\text{off}} = |I^+ - I^-|$



Le gain en boucle ouverte ainsi que la bande passante ne sont pas infinis. En dehors de la bande passante, l'amplification décroît de 20 dB par décade.



On peut déduire de la remarque précédente que la pente maximale de sortie (MSR pour *Max Slew Rate*) est également limitée, c'est-à-dire que le signal en sortie de l'AOP réel a une vitesse de variation limitée :

$$\text{MSR} = \left. \frac{dv_s(t)}{dt} \right|_{\text{max}}$$

49 Montages à amplificateur opérationnel

Mots clés

Montages non inverseur, inverseur, amplificateur différentiel, suiveur, sommateur, intégrateur, filtres actifs.

1. EN QUELQUES MOTS

Dans cette fiche, nous considérerons que les amplificateurs opérationnels sont idéaux (impédance d'entrée infinie) :

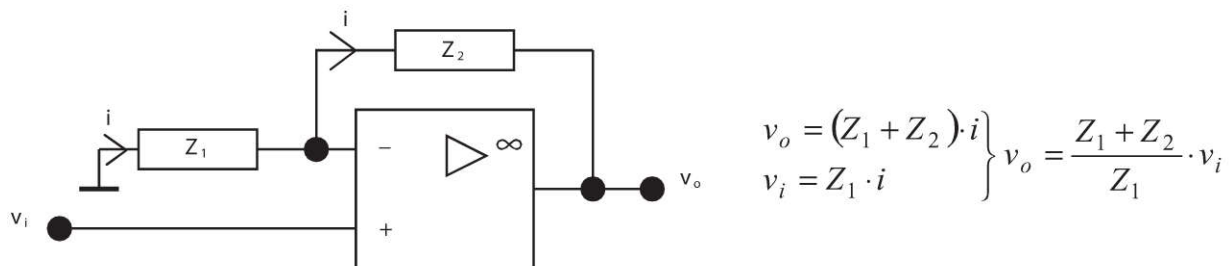
- les courants d'entrée sont nuls ;
- la différence de potentiel à l'entrée est nulle ;
- l'amplification est infinie.

Les montages proposés ici sont autant de fonctions très répandues en électronique.

2. CE QU'IL FAUT RETENIR

a) Montage non inverseur

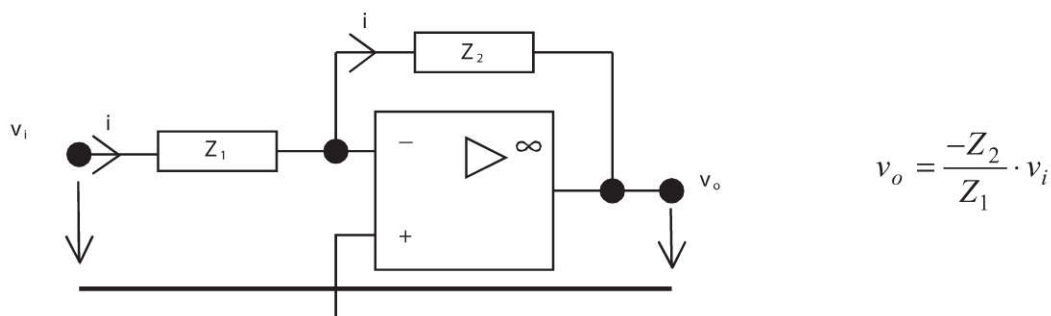
L'impédance d'entrée de l'AOP étant infinie, la totalité du courant i traversant Z_1 traverse Z_2 .



Avec Z_1 et Z_2 des résistances, ce montage est un **amplificateur de tension** dont le gain dépend du choix des résistances.

b) Inverseur

La tension de sortie est inversée par rapport à l'entrée : l'entrée non inverseuse est reliée à la masse. De même que précédemment, la totalité du courant i traversant Z_1 traverse Z_2 .

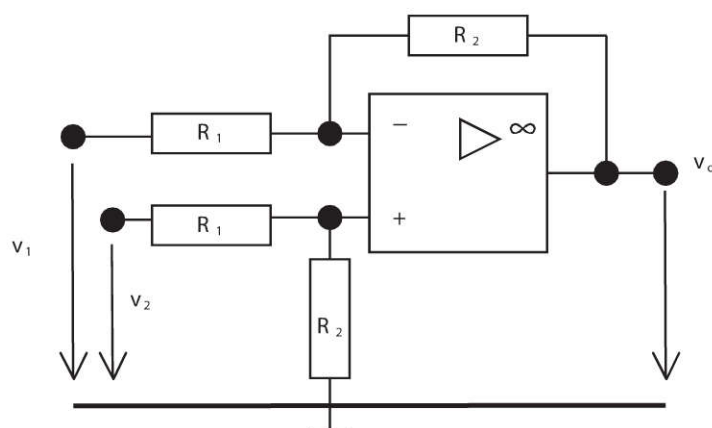


L'impédance d'entrée s'écrit : $Z_{in} = Z_1$

Lorsque ce montage n'est pas utilisé en filtre ([fiche 50](#)), Z_1 et Z_2 sont des résistances. Le gain de ce montage peut être supérieur à 1 et le choix des valeurs de résistances permet de **régler ce gain**.

c) Amplificateur différentiel

En combinant les montages précédents et en appliquant le théorème de superposition, l'expression de la tension de sortie du montage est fonction de la différence des tensions d'entrée et avec un facteur de gain dépendant des valeurs de résistances.



$$v_o = (v_2 - v_1) \cdot \frac{R_2}{R_1}$$

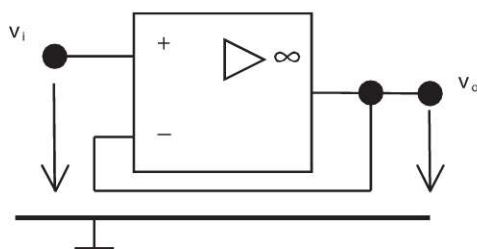
$$Z_{in}|_{v_1} = R_1$$

$$Z_{in}|_{v_2} = R_1 + R_2$$

Cette expression est à rapprocher du fonctionnement d'un amplificateur différentiel (fiche 48). Ce montage est également appelé **soustracteur**.

d) Suiveur

Le suiveur de tension est un montage d'**adaptation d'impédance**, c'est-à-dire que l'on évite d'avoir un montage diviseur de tension entre l'impédance de la source et celle de la charge : c'est ce que l'on appelle le découplage.

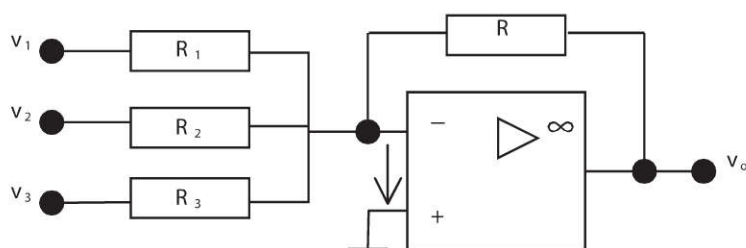


$$v_o = v_i$$

Idéalement, le gain est égal à 1 (le montage n'amplifie pas et n'atténue pas le signal), l'impédance d'entrée est infinie, l'impédance de sortie est nulle. Ce montage est appelé tampon ou buffer car il permet de piloter une charge avec un signal ayant un faible support de charge.

e) Sommateur

Le montage sommateur est une évolution du montage inverseur :



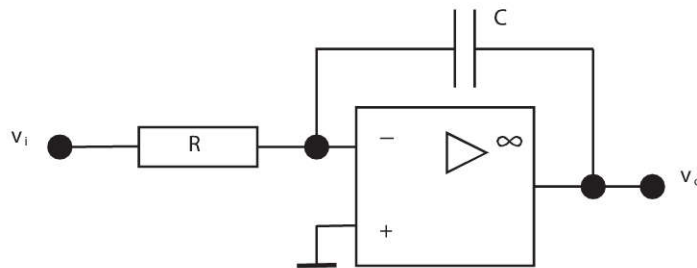
$$v_o = -R \cdot \left(\frac{v_1}{R_1} + \frac{v_2}{R_2} + \frac{v_3}{R_3} \right)$$

$$= -R \cdot \sum_i \frac{v_i}{R_i}$$

$$v_o = -\sum_i v_i \text{ si les } R_i \text{ sont identiques et égales à } R.$$

f) Intégrateur

En remplaçant Z_1 par R et Z_2 par une capacité C telle que $i(t) = C \cdot dv_C(t)/dt$ dans le montage inverseur, on trouve :



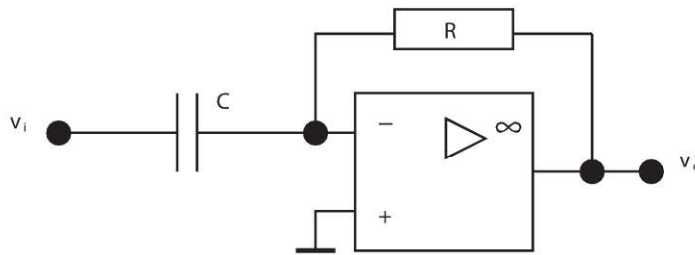
$$v_o = -\frac{1}{R \cdot C} \cdot \int_{t_0}^t v_i(t) dt$$

avec $v_i(0) = 0$

Ce montage est appelé intégrateur du fait de la forme de sa fonction de transfert. Il s'apparente à un filtre d'ordre 1 (fiche 50).

g) Dérivateur ou différentiateur

En utilisant la formule de l'impédance de l'inductance et en remplaçant Z_1 par $j \cdot L \cdot \omega$ et Z_2 par R dans le montage inverseur, on trouve :

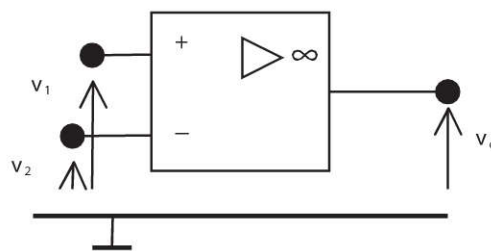


$$v_o = -R \cdot C \cdot \frac{dv_i(t)}{dt}$$

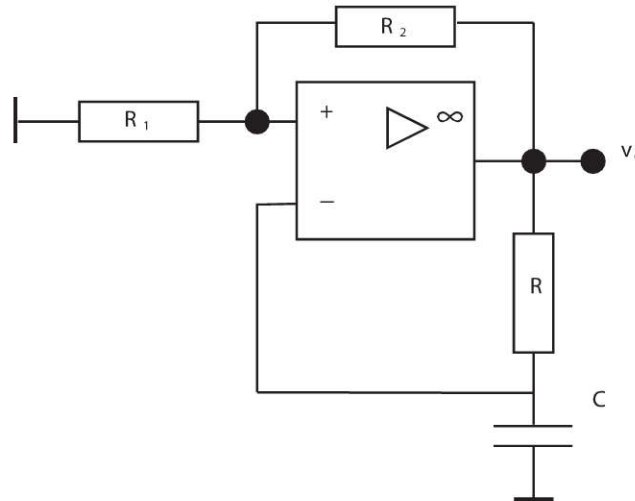
Ce montage est appelé différentiateur du fait de la forme de sa fonction de transfert. Il s'apparente également à un filtre d'ordre 1 (fiche 50).

h) Comparateur

Lorsque la tension v_1 est supérieure à v_2 , la tension en sortie v_o est maximum (tension d'alimentation). Dans le cas contraire ($v_1 < v_2$), la tension de sortie est minimum (0 V).



i) La bascule astable ou multivibrateur astable



Le bouclage sur l'entrée + entraîne un fonctionnement non linéaire de l'amplificateur opérationnel, la sortie va donc basculer entre deux valeurs de tensions qui correspondent aux tensions de saturation de l'AOP. Celles-ci sont, en première approximation, égales aux tensions d'alimentation du circuit.

La période d'oscillation est $T = 2 \cdot R \cdot C \cdot \ln \left(1 + 2 \frac{R_1}{R_2} \right)$

3. EN PRATIQUE

a) Filtres actifs

Les filtres actifs à base d'amplificateurs opérationnels font l'objet des [fiches 50](#) et [51](#). Notons que les filtres actifs d'ordre 1 se basent sur les montages inverseur et non-inverseur vus dans cette fiche.

b) CNA

Les convertisseurs numérique-analogique (CNA, [fiche 62](#)) utilisent le montage sommateur (CNA R-2R ou à résistance pondérées).

c) CAN

Certains convertisseurs analogique-numérique (CAN, [fiche 63](#)) utilisent le montage comparateur, comme typiquement dans les convertisseurs flash.

50 Filtres actifs d'ordre 1 et 2

Mots clés

Cellule de Rauch, de Sallen-Key.

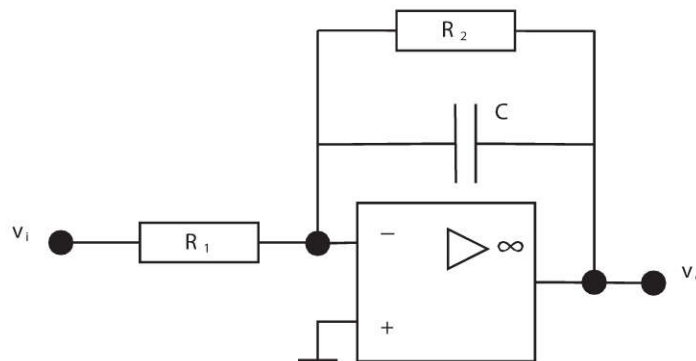
1. EN QUELQUES MOTS

Un filtre actif est un circuit constitué d'un amplificateur opérationnel et de composants passifs. Contrairement aux filtres passifs, ils peuvent éventuellement amplifier le signal. Son intérêt est de favoriser l'intégration des systèmes (compacité) en évitant l'utilisation d'inductances souvent volumineuses et peu précises : les cellules de filtrages présentées utilisent donc uniquement des condensateurs et des résistances.

2. CE QU'IL FAUT RETENIR

a) 1^{er} ordre

Une cellule de filtrage de premier ordre, c'est-à-dire dont la fonction de transfert est une équation différentielle d'ordre 1, n'emploie qu'une seule résistance. Le schéma proposé est un exemple de cellule de filtrage passe-bas du premier ordre.

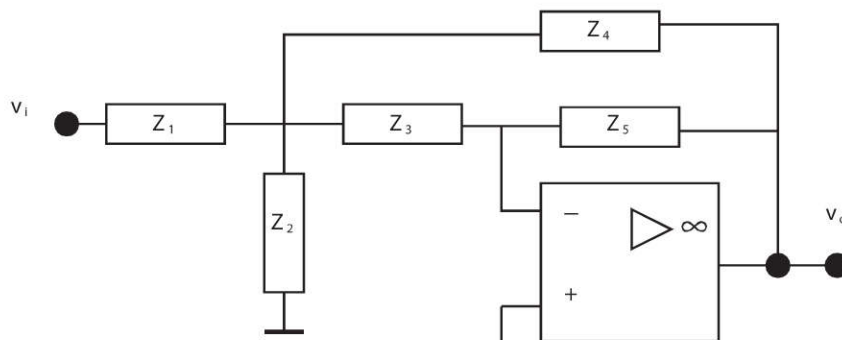


Sa fonction de transfert, après transformée de Laplace, est :

$$\frac{U_o(p)}{U_i(p)} = \frac{R_2/R_1}{R_2 \cdot C \cdot p + 1}$$

b) Rauch

Ces cellules sont utilisées pour synthétiser des filtres d'ordre 2 de type passe-bande.



En toute généralité, sa fonction de transfert s'écrit :

$$v_s = \frac{-v_e}{\frac{Z_3}{Z_5} + \frac{Z_1}{Z_5} + \frac{Z_1}{Z_4} + \frac{Z_1 \cdot Z_3}{Z_2 \cdot Z_5} + \frac{Z_1 \cdot Z_3}{Z_4 \cdot Z_5}}$$

Soient :

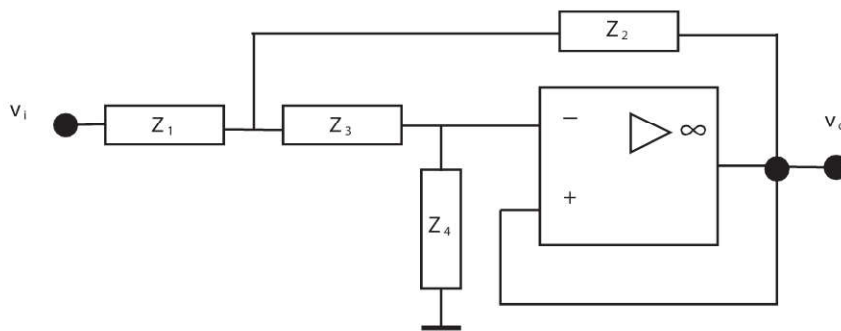
- $Z_1 = R_1$
- $Z_2 = R_2$
- $Z_3 = Z_4 = C$
- $Z_5 = R_3$

Nous obtenons :

$$H(p) = \frac{-k \cdot R \cdot C \cdot p}{k \cdot R_1 \cdot R_3 \cdot C^2 \cdot p^2 - 2 \cdot k \cdot R_1 \cdot C \cdot p + 1} \quad \text{avec } k = \frac{R_2}{R_1 + R_2} \text{ (diviseur de tension)}$$

c) Sallen-Key

Ces cellules sont utilisées pour synthétiser des filtres d'ordre 2 de type passe-haut ou passe-bas et se présentent sous cette forme :



Soient :

- $Z_1 = Z_3 = R$
- $Z_2 = C_1$
- $Z_4 = C_2$

Voici les expressions de la fonction de transfert $H(p)$ (domaine de Laplace), du coefficient de qualité Q et de la pulsation de résonance ω_0 .

$$H(p) = \frac{1}{R_2 \cdot C_1 \cdot C_2 \cdot p^2 + 2 \cdot R \cdot C_2 + 1}$$

$$Q = \frac{1}{2} \sqrt{\frac{C_1}{C_2}}$$

$$\omega_0 = \frac{1}{R \cdot \sqrt{C_1 \cdot C_2}}$$

3. EN PRATIQUE

À partir des fonctions de transfert, il est possible de dimensionner les valeurs des différents composants passifs. La mise en œuvre de ces cellules de filtrage fait l'objet de la fiche suivante (fiche 52).

51 Synthèse de filtres actifs

Mots clés

Gabarit, sélectivité, symétrie, normalisation, dénormalisation, spectre.

1. EN QUELQUES MOTS

Modifier les composantes spectrales d'un signal de façon à le faire « rentrer » dans un gabarit. On désigne par « gabarit » la réponse de l'amplification du signal en fonction de la fréquence, en imposant des valeurs spécifiques. On note $X(f)$ la transformée de Fourier du signal d'entrée du filtre et $Y(f)$ sa sortie.

2. CE QU'IL FAUT RETENIR

a) Caractéristiques d'un filtre : fonction de transfert

Le filtre est généralement un système linéaire invariant, il est donc entièrement caractérisé par sa fonction de transfert (ou sa réponse impulsionnelle complexe). De ce fait, les propriétés que l'on souhaite voir réaliser par le filtre sont imposées directement sur sa fonction de transfert par l'intermédiaire de gabarits (phase et amplitude). Sa réalisation peut se faire avec des éléments passifs (R, L, C) ou actifs (ALI fonctionnant dans les structures de Rauch et de Salley-Key)

b) Spécification des caractéristiques d'un filtre

Un **gabarit** permet de spécifier le comportement spectral d'un filtre.

- En module par $H_{dB} = 10 \cdot \log_{10} \left(\left| \frac{Y(f)}{X(f)} \right|^2 \right)$ c'est-à-dire par le module de sa fonction de

transfert, mais bien souvent par son atténuation : $A_{dB} = 10 \cdot \log_{10} \left(\left| \frac{X(f)}{Y(f)} \right|^2 \right) = -H_{dB}$

- En phase, en imposant l'évolution du déphasage entre $Y(f)$ et $X(f)$

Gabarits et fréquence de référence

On définit plusieurs gabarits en fonction de la nature du filtre. Les paramètres sont :

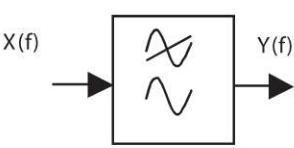
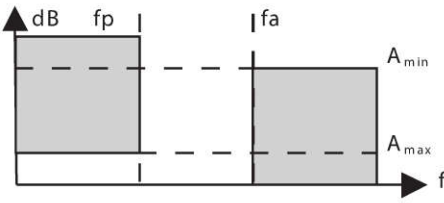
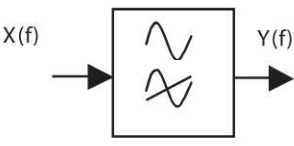
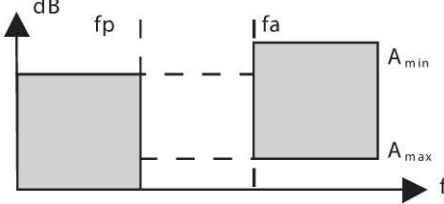
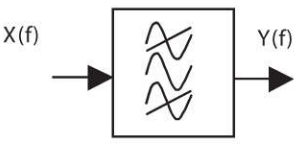
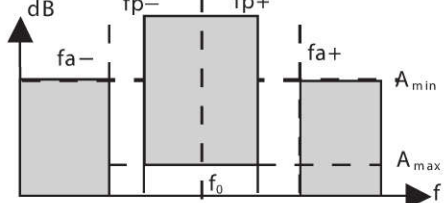
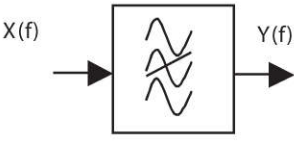
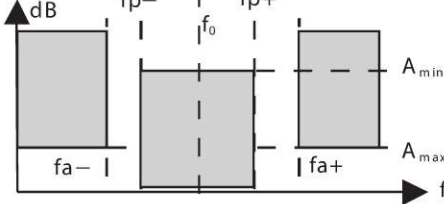
- **Amax** = l'atténuation maximale dans la bande passante exprimée en dB ;
- **Amin** = l'atténuation minimale dans la bande atténuée exprimée aussi en dB ;
- les fréquences caractérisant les zones de transition entre la bande passante (f_{p+} et f_{p-}) et la bande atténuée (f_{a+} et f_{a-}) ;
- la sélectivité K, ou la raideur de la bande de transition, est définie à partir des fréquences de coupures.

Sélectivité	Passe-bas	Passe-haut	Passe-bande	Coupe-bande
K	$\frac{f_p}{f_a}$	$\frac{f_a}{f_p}$	$\frac{\Delta f_p}{\Delta f_a} = \frac{f_{p+} - f_{p-}}{f_{a+} - f_{a-}}$	$\frac{\Delta f_a}{\Delta f_p} = \frac{f_{a+} - f_{a-}}{f_{p+} - f_{p-}}$

Pour les filtres passe-bande ou coupe-bande symétriques, on peut définir la fréquence centrale du filtre symétrique : $f_0 = f_a^+ \propto f_a^- = f_p^+ \propto f_p^-$ (moyenne logarithmique de deux grandeurs) et la largeur de bande relative B telles que :

Passe-bande	Coupe-bande
$B = \frac{\Delta f_p}{\Delta f_0} = \frac{f_p^+ - f_p^-}{f_0}$	$B = \frac{\Delta f_a}{\Delta f_0} = \frac{f_a^+ - f_a^-}{f_0}$

Les gabarits sont les suivants :

Type de filtre	Représentation schématique	Gabarit en atténuation	Fréquence de référence
Passe-bas			$f_n = f_p$
Passe-haut			$f_n = f_p$
Passe-bande			$f_n = f_0 = \sqrt{f_p^+ \times f_p^-}$
Coupe-bande			$f_n = f_0 = \sqrt{f_p^+ \times f_p^-}$

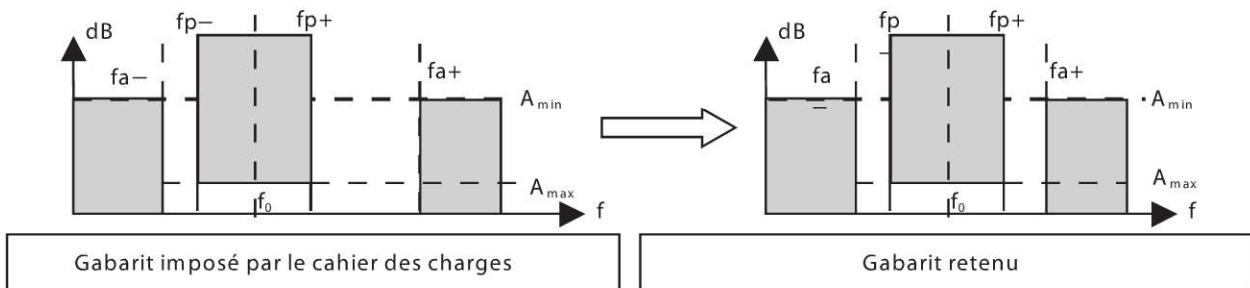
On note la fréquence normalisée la fréquence $F = \frac{f}{f_n}$. Ainsi, tous les filtres dont les courbes de réponse sont semblables à un rapport de fréquence près, auront la même représentation.

c) Symétrisation des filtres coupe et passe-bande

Pour des raisons pratiques dues aux méthodes de synthèse de filtre, il est difficile de synthétiser un filtre dont le gabarit n'est pas symétrique par rapport à la fréquence de référence. On utilise alors un gabarit symétrique qui est inclus dans un gabarit initial. On synthétise donc un filtre pré-

sentant un **gabarit** qui est **plus contraignant** que le gabarit d'origine. Le filtre réalisé sera donc en accord avec le cahier des charges.

Exemple pour un filtre passe-bande



Qui peut le plus, peut le moins... la bande de fréquence $[fp^+, fa^+]$ est plus contraignante pour le gabarit retenu que le gabarit imposé par le cahier des charges, mais cela rendra son étude plus simple. Même si la contrainte sur les composants permettant de le réaliser sera tout aussi importante.

3. EN PRATIQUE

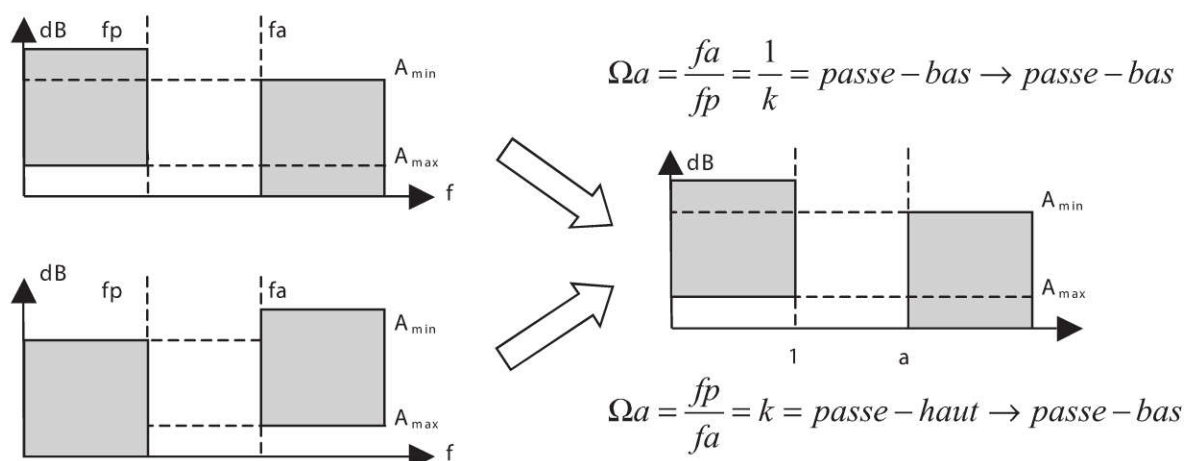
a) Prototypage sur la base d'un filtre passe bas

Même en utilisant des unités normalisées, le nombre de cas à prendre en considération dans l'étude du filtrage est considérable. C'est pourquoi on a l'habitude de n'étudier que les filtres passe-bas et de ramener la réalisation des filtres passe-haut, passe-bande et coupe-bande à celle d'un filtre **passe bas normalisé** appelé **prototype**.

Le prototypage entraîne une modification de la variable à manipuler. Ainsi, la variable de Laplace p devient P , la variable de Laplace normalisée.

La fonction de transfert du filtre avant normalisation $H(p) = H(j)$ devient $H(P) = H(j)$
 $= H(j2F)$ après normalisation.

Ainsi un filtre passe-bas et un filtre passe-haut lorsqu'ils deviennent normalisés, donnent :



Les amplitudes ne sont pas affectées par la normalisation.

b) Transposition

Quel que soit le gabarit du filtre à réaliser (ce dernier devant être symétrique dans le cas des filtres passe-bande ou coupe-bande), on se ramène d'abord au gabarit du filtre prototype par une

transposition conforme. Une fois la fonction de transfert de ce filtre connue, on revient au filtre désiré par la même transposition.

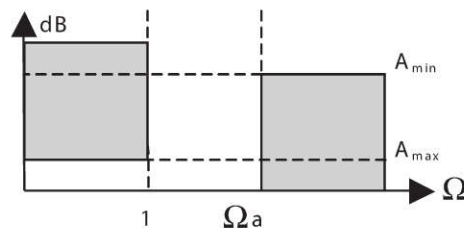
On remplace le P de l'expression du filtre passe-haut normalisé par $(1/P)$ pour obtenir l'expression du filtre passe-bas prototype.

On « dénormalise » le filtre passe-bas prototype en remplaçant le P de l'expression du filtre par $(1/P)$ pour obtenir l'expression du filtre passe-haut normalisé. Le tableau ci-dessous résume les trois cas de transposition.

Filtre d'origine	Transposition	Filtre obtenu
passe-bas	$P1/P$	passe-haut
passe-bas	$P \leftrightarrow \frac{1}{B} \times \left(P + \frac{1}{P} \right)$	passe-bande
passe-bas	$P \leftrightarrow \frac{1}{\frac{1}{B} \times \left(P + \frac{1}{P} \right)}$	coupe-bande

c) Polynômes

Le gabarit d'un filtre passe-bas normalisé apparaît sous la forme ci-dessous. On suppose que la fréquence de normalisation est égale à la fréquence de la bande passante (f_p) correspondant à l'atténuation $A_{\max}$. La fréquence correspondant au début de la bande atténuée est f_a et la valeur normalisée correspondante est a .



Toute la subtilité est de trouver la fonction mathématique qui permet d'inscrire une courbe dans ce gabarit. C'est-à-dire qui soit une approximation de la courbe idéale qui est :

- une atténuation faible dans la bande passante ($f < f_p$) ;
- une atténuation forte dans la bande atténuée ($f > f_a$) ;
- une atténuation qui doit augmenter rapidement entre f_p et f_a .

On note l'atténuation relative en puissance $A^2(\Omega) = 1 + \epsilon^2 k_n^2(\Omega)$

$K_n(\Omega)$ est une **fonction d'approximation** telle que $K_n(1) = 1$

Si $H(\Omega)$ est d'ordre n alors $A(\Omega)$ est d'ordre $2n$ et contient $2n$ pôles.

On choisit les pôles stables pour $H(P)$, c'est-à-dire à partie réelle négative de $1 + \epsilon^2 k_n^2(-jP) = 0$

Les inconnues sont l'ordre du filtre (n), le coefficient et la fonction $K_n()$ qui dépend du comportement voulu, dans et à l'extérieur, de la bande.

52 Logique booléenne

Mots clés

Variable logique, binaire, équation logique, opérateur ET, OU, NON, XOR, NAND, NOR, tableau de Karnaugh, théorèmes de De Morgan.

1. EN QUELQUES MOTS

Une **variable binaire** (ou variable logique) ne peut prendre ses valeurs que dans l'ensemble $\{0, 1\}$. Un 1 ou un 0 peut par exemple correspondre à la position d'un interrupteur ou à un niveau de tension. Une **équation** logique est le résultat de la combinaison de plusieurs variables logiques. La **logique booléenne** est un **formalisme mathématique** décrivant ces relations qui est à la base de la théorie des systèmes électroniques numériques.

2. CE QU'IL FAUT RETENIR

a) Symboles

La logique booléenne emploie très peu de symboles différents :

- $+$ représente l'opérateur logique OU (somme logique) ;
- $\cdot$ représente l'opérateur logique ET (produit logique) ;
- $\oplus$ représente l'opérateur logique OU EXCLUSIF ;
- $\overline{A}$ représente le complément de A (inversion).

Ces opérateurs correspondent à des règles très simples :

- $A + B = 1$ si A ou B est égal à 1 ($A + B = 0$ dans le cas contraire) ;
- $A \cdot B = 1$ si A et B sont simultanément égaux à 1 ($A \cdot B = 0$ dans le cas contraire) ;
- $A \oplus B = 1$ si A et B sont différents ($A \oplus B = 0$ dans le cas contraire) ;
- $\overline{A} = 1$ si A vaut 0 et $\overline{A} = 0$ si A vaut 1

Cela peut se traduire par les **tables de vérité** suivantes :

A	$\overline{A}$
0	1
1	0

B	A	$A \cdot B$
0	0	0
0	1	0
1	0	0
1	1	1

B	A	$A + B$
0	0	0
0	1	1
1	0	1
1	1	1

B	A	$A \oplus B$
0	0	0
0	1	1
1	0	1
1	1	0

Ces opérateurs peuvent être combinés afin de constituer des équations logiques (de façon classique, les parenthèses indiquent la priorité).

b) Propriétés

► **Commutativité**

L'opérateur OU est commutatif : $A + B = B + A$

L'opérateur ET est commutatif : $A \cdot B = B \cdot A$

► **Associativité**

L'opérateur OU est associatif : $A + (B + C) = (A + B) + C = A + B + C$

L'opérateur ET est associatif : $A \cdot (B \cdot C) = (A \cdot B) \cdot C = A \cdot B \cdot C$

► **Distributivité**

L'opérateur OU est distributif sur l'opérateur ET : $(A \cdot B) + C = A + C \cdot B + C$

L'opérateur ET est distributif sur l'opérateur OU : $(A + B) \cdot C = A \cdot C + B \cdot C$

► **Élément neutre**

1 est l'élément neutre du produit logique : $A \cdot 1 = A$

0 est l'élément neutre de la somme logique : $A + 0 = A$

► **Élément absorbant**

1 est l'élément absorbant de la somme logique : $A + 1 = 1$

0 est l'élément absorbant du produit logique : $A \cdot 0 = 0$

► **Opérations impliquant deux fois la même variable**

$$A + A = A \qquad A + \bar{A} = 1 \qquad A \cdot A = A \qquad A \cdot \bar{A} = 0$$

c) Opérateurs logiques NOR et NAND

L'opération combinant la complémentation et l'opérateur OU est appelée NON-OU ou plus communément NOR (not-or) : $\overline{A + B}$

L'opération combinant la complémentation et l'opérateur ET est appelée NON-ET ou plus communément NAND (not-and) : $\overline{A \cdot B}$

Ces deux opérations possèdent la propriété d'**idempotence**, c'est-à-dire que chacune d'elles permet d'obtenir toutes les autres. Les opérateurs NAND et NOR sont dits **universels**.

d) Théorèmes de De Morgan

• 1^{er} théorème : $\overline{A + B}$

• 2^e théorème : $\overline{A \cdot B}$

Soient Σ l'opération de somme logique sur n variables logiques A_i et Π l'opération de produit logique sur n variables logiques A_i . Les théorèmes de De Morgan peuvent donc se généraliser sur n variables de la manière suivante :

• 1^{er} théorème : $\overline{\prod_{i=1}^N A_i} = \sum_{i=1}^N \bar{A}_i$

• 2^e théorème : $\overline{\sum_{i=1}^N A_i} = \prod_{i=1}^N \bar{A}_i$

3. EN PRATIQUE

La **méthode de Karnaugh** permet de simplifier une fonction ou une expression logique. Elle est basée sur un tableau de $2n$ cases, où n est le nombre de variables. Les entrées du tableau indiquent l'état des variables d'entrée codées en code Gray (**fiche 53**), et dans chaque case, on place l'état de la sortie pour les combinaisons d'entrée correspondantes.

Voici un exemple à quatre variables A, B, C et D.

A, B \ C, D	00	01	10	11
00	0	0	0	0
01	1	0	0	1
11	1	1	1	1
10	0	0	0	0

Après avoir rempli ce tableau, il faut effectuer des groupements d'états de sortie à 1 dans des cases adjacentes ou voisines. Pour cela, il faut suivre les règles suivantes afin de trouver une équation logique simplifiée :

- regrouper les 1 par nombre d'une puissance de 2 (2, 4, 8, 16...);
- un 1 peut faire partie de plusieurs ensembles : il faut constituer les ensembles les plus grands possibles, quitte à ce que des éléments soient inclus dans plusieurs ensembles en même temps ;
- les « bords » du tableau sont des côtés adjacents ;
- les regroupements en diagonale ne sont pas permis.

Voici la meilleure possibilité de regroupement pour l'exemple proposé :

A, B \ C, D	00	01	11	10
00	0	0	0	0
01	1	0	0	1
11	1	1	1	1
10	0	0	0	0

Voici comment obtenir l'équation logique avec cet exemple :

- pour le regroupement en rouge : $B = 1$ et $D = 0$ quels que soient A et C ;
- pour le regroupement en gris : $A = B = 1$ quels que soient C et D.

La sortie est à 1 si [A ET B sont à 1] OU si [B est à 1 ET D à 0]

L'équation logique peut donc s'écrire $A \cdot B + B \cdot \overline{D}$

Voici quelques exemples de regroupements possibles :

A, B \ C, D	00	01	11	10
00	0	1	1	0
01	1	0	0	1
11	1	0	0	1
10	0	1	1	0

$$S = \overline{B} \cdot D + B \cdot \overline{D}$$

A, B \ C, D	00	01	11	10
00	1	0	0	1
01	0	0	0	0
11	0	0	0	0
10	1	0	0	1

$$S = \overline{B} \cdot \overline{D}$$

A, B \ C, D	00	01	11	10
00	0	1	0	0
01	1	0	0	0
11	0	0	1	0
10	0	0	0	0

$$S = \overline{A} \cdot B \cdot \overline{C} \cdot \overline{D} + \overline{A} \cdot \overline{B} \cdot \overline{C} \cdot D + A \cdot B \cdot C \cdot D \quad S = B$$

A, B \ C, D	00	01	11	10
00	0	0	0	0
01	1	1	1	1
11	1	1	1	1
10	0	0	0	0

Voici quelques exemples de regroupements impossibles au sens de la méthode de Karnaugh :

A, B \ C, D	00	01	11	10
00	0	1	0	0
01	1	0	0	0
11	0	0	0	0
10	1	0	1	0

A, B \ C, D	00	01	11	10
00	0	1	1	1
01	1	1	1	0
11	1	1	1	0
10	0	0	0	0

53 Codes numériques

Mots clés

Complément à 2, binaire décalé, BCD, DCB, hexadécimal, octal.

1. EN QUELQUES MOTS

Les variables binaires ne permettent de coder que des 1 ou des 0. En électronique, ces variables sont appelées **bit**, pour binary digit (nombre binaire). Partant de là, il est possible de **regrouper plusieurs bits** pour former des **mots** (ou **nombres**) de plusieurs bits : différents codes numériques existent pour établir une relation entre un mot binaire et un autre code.

2. CE QU'IL FAUT RETENIR

a) Principe décimal et code binaire naturel

De même qu'en décimal, les chiffres de droite valent moins que les chiffres de gauche : dans « 1 427 », le 1 vaut plus que le 4 qui vaut plus que le 2 qui lui-même vaut plus que le 7. 1 427 peut se décomposer en :

$$1 \times 1\,000 + 4 \times 100 + 2 \times 10 + 7 \times 1 = 1\,427.$$

Le rang du chiffre indique le nombre de zéros à rajouter derrière ce chiffre. 1 a le 4^e rang, celui des milliers → 3 zéros. 4 a le 3^e rang, celui des centaines → 2 zéros. 2 a le 2^e rang, celui des dizaines → 1 zéro. On pourrait aussi écrire :

$$1 \times 10 \times 10 \times 10 + 4 \times 10 \times 10 + 2 \times 10 + 7 \times 1 = 1\,427.$$

Le rang représente donc le nombre de fois où l'on multiplie 10 par lui-même. L'écriture avec des puissances de 10 donnerait :

$$1 \times 10^3 + 4 \times 10^2 + 2 \times 10^1 + 7 \times 10^0 = 1\,427.$$

Traduisons $(10011101)_2$ en décimal. Ici, le 1^{er} rang vaut 1, le 2^e vaut 2, le 3^e vaut 4, le 4^e vaut 8, etc. Ici, au lieu de puissances de 10, nous manipulons des puissances de 2 ! On a donc : $(10011101)_2 =$

$$1 \times 2^7 + 0 \times 2^6 + 0 \times 2^5 + 1 \times 2^4 + 1 \times 2^3 + 1 \times 2^2 + 0 \times 2^1 + 1 \times 2^0$$

$$\text{Soit } (10011101)_2 = 1 \times 128 + 0 \times 64 + 0 \times 32 + 1 \times 16 + 1 \times 8 + 1 \times 4 + 0 \times 2 + 1 \times 1$$

$$\text{Donc } (10011101)_2 = 128 + 16 + 8 + 4 + 1 = (157)_{10}$$

Cette dernière ligne peut se lire « un, zéro, zéro, un, un, un, zéro, un, en base deux est égal à cent cinquante-sept en base dix ».

Ce principe permet donc de retrouver un nombre décimal entier positif (entier naturel) à partir du code **binaire naturel**, c'est le **binaire codé décimal ou DCB**.

L'opération inverse, le **décimal codé binaire ou DCB**, implique de retrouver la puissance de 2 immédiatement inférieure au nombre décimal entier, de retrancher la valeur correspondante et d'effectuer la même opération avec les restes jusqu'à obtenir 0.

2^0	2^1	2^2	2^3	2^4	2^5	2^6	2^7	2^8	2^9	2^{10}	2^{11}	2^{12}
1	2	4	8	16	32	64	128	256	512	1 024	2 048	4 096

Prenons 87 comme exemple. La puissance de 2 immédiatement inférieure est 64 (2^6) : le nombre binaire final devra donc comporter 7 bits.

$87 - 2^6 = 87 - 64 = 23$	1	Bit de poids fort
$23 - 2^5 = 23 - 32 < 0$	0	
$23 - 2^4 = 23 - 16 = 7$	1	
$7 - 2^3 = 7 - 8 < 0$	0	
$7 - 2^2 = 7 - 4 = 3$	1	
$3 - 2^1 = 3 - 2 = 1$	1	
$1 - 2^0 = 1 - 1 = 0$	1	Bit de poids faible

Nous obtenons donc $(87)_{10} = (1010111)_2$.

Le repérage du **bit de poids fort** (ou MSB pour *Most Significant Bit*) et du **bit de poids faible** (ou LSB pour *Least Significant Bit*) indique dans quel sens lire le nombre binaire.

b) Nombres entiers négatifs

Le codage binaire de nombres entiers négatifs implique d'ajouter un **bit de signe** au nombre binaire. La solution la plus simple consisterait à rajouter ce bit au code binaire naturel :

Décimal	BCD avec bit de signe
3	011
2	010
1	001
0	000
(-0)	100
-1	101
-2	110
-3	111

Deux inconvénients s'opposent à l'utilisation de cette solution intuitive : d'une part, le 0 décimal est codé deux fois, et d'autre part, cette technique n'est pas compatible avec l'addition binaire. Deux techniques permettent de palier à ces problèmes.

Le **code binaire décalé** consiste, comme son nom l'indique, à décaler le code binaire naturel « vers le bas » de sorte à obtenir des nombres négatifs. Le premier bit est le bit de signe : il est à 0 pour les nombres négatifs et à 1 pour les nombres positifs.

Le **code binaire complément à 2**. Le premier bit est le bit de signe : il est à 1 pour les nombres négatifs et à 0 pour les nombres positifs. Pour obtenir ce code, il faut :

- inverser les bits du BCD (complémentation), ce qui donne un code complément à 1 ;
- ajouter 1 aux nombres obtenus.

Entre ces deux codes, le **bit de signe est complémenté** (inversé).

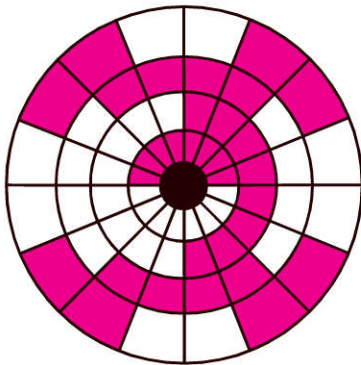
Décimal	Binaire décalé	Complément à 2
-1	011	111
-2	010	110
-3	001	101
-4	000	100

Décimal	Binaire décalé	Complément à 2
3	111	011
2	110	010
1	101	001
0	100	000

c) Code de Gray

Le code de Gray, ou **code binaire réfléchi**, n'est pas spécifiquement conçu pour coder des nombres. Il est basé sur le principe de ne changer qu'un seul bit entre deux nombres binaire successifs (ce qui n'est pas le cas en binaire naturel : entre 3 et 4 en base 10, les trois bits changent en binaire naturel).

Ce code permet d'éviter les comportements aléatoires de certains systèmes. Par exemple, dans les souris d'ordinateurs munies d'une boule, les déplacements transversaux et longitudinaux sont chacun associés à une roue codeuse : un petit déplacement sur un des axes fait tourner la roue (percée de trous selon le code de Gray) et le plus petit pas de rotation possible est associé au changement « d'état » d'un seul trou.



Décimal	Binaire naturel	Binaire réfléchi
7	111	100
6	110	101
5	101	111
4	100	110
3	011	010
2	010	011
1	001	001
0	000	000

Pour trouver un code à partir du précédent, il faut inverser le bit le plus à droite possible permettant d'obtenir un nouveau nombre. De façon pratique, le nom de code binaire réfléchi provient d'une technique « miroir » (d'où le nom « réfléchi ») permettant d'obtenir un code de Gray sur N bits.

1. Partons des deux nombres 0 et 1.
2. Pour chaque bit supplémentaire, on recopie en miroir les nombres de l'étape précédente.
3. Il faut rajouter un 0 au début des nombres issus de l'étape précédente, et un 1 aux nombres obtenus par symétrisation.

Gray 1 bit	Gray 2 bits	Gray 3 bits
0	0 0	0 00
		0 01
	0 1	0 10
		0 11
1	1 1	1 11
		1 10
	1 0	1 01
		1 00

d) Nombre décimaux

La représentation en **virgule fixe** se base sur le code complément à deux. Les chiffres binaires placés avant la virgule sont des entiers puissance de 2 et les nombres placés après valent respectivement $1/2$, $1/4$, $1/8$... soit $1/2^N$ avec N le rang du chiffre binaire. Traduisons $(1011,101)_2$ en décimal.

Binaire	1	0	1	1		1	0	1
Décimal	8	4	2	1	,	0,5	0,25	0,125
Nombre virgule fixe	8	0	2	1		0,5	0	0,125

Nous obtenons $(1011,101)_2 = 8 + 2 + 1 + 0,5 + 0,125 = (9,625)_{10}$

La représentation en **virgule flottante** se base sur l'expression suivante :

$$r = S_M \cdot M \cdot 10^{S_E \cdot E}$$

r : réel à coder

S_M : signe de la mantisse (1 négatif, 0 positif)

S_E : signe de l'exposant (1 négatif, 0 positif)

M : mantisse, elle représente les chiffres significatifs en binaire naturel

E : exposant en binaire naturel

Cette représentation utilise 16, 32, 64 ou 128 bits en fonction des processeurs et des compilateurs (fiche 56).

Nombre de bits	Exposant N	Mantisse M
16	4	10
32	8	22
64	14	48

Exemple

À quel réel correspond des nombres binaires virgule flottante 101110110100000 et 0100110000101011 ?

S_M	S_E	E	M		S_M	S_E	E	M	
1	0	1111	0110100000		0	1	0011	0000101011	
-	+	15	224	$= -224 \cdot 10^{15}$	+	-	3	43	$= -43 \cdot 10^{-3}$

3. EN PRATIQUE

La capacité des mémoires est indiquée en octets, c'est-à-dire en groupant les bits par 8 (8 bits = un octet). On utilise souvent le mot anglais **byte** pour parler d'octets. Le bit se note b et le byte se note B.

Une clé USB de 1 Go, soit un Gigaoctet ou 1 GB (Gigabyte), peut stocker 8 Gigabits.

Cette clé ne comporte pas exactement 1 000 000 000 d'octets (un milliard), mais un peu plus.

Avec 10 bits on peut représenter 1 024 valeurs ($= 2^{10}$), c'est-à-dire $2 \times 2 \times 2 \times 2 \times 2 \times 2 \times 2 \times 2 \times 2 \times 2$ valeurs. En informatique et en électronique numérique, on associe le mot kilo au nombre 1 024. Méga est associé à $1\,024 \times 1\,024$ soit 1 048 576. Le Giga de notre clé USB, de la même manière, ne représente pas 1 milliard, mais 1 024 Méga soit 1 073 741 824.

54 Portes logiques, logique combinatoire

Mots clés

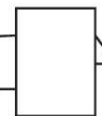
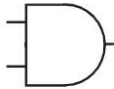
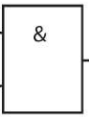
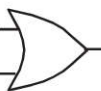
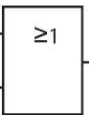
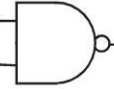
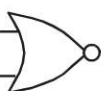
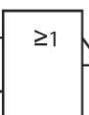
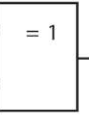
Logique booléenne, circuit intégré, théorème de De Morgan.

1. EN QUELQUES MOTS

Les portes logiques sont des composants numériques qui fonctionnent avec des signaux logiques à deux états (0 ou 1). Ils sont utilisés dans les domaines de l'électronique numérique, de l'automatique ou encore de l'électronique de puissance pour des commandes automatiques d'interrupteurs.

2. CE QU'IL FAUT RETENIR

Portes logiques

Fonction logique	Symbole		Opération booléenne	Table de vérité															
	américain	européen																	
NON « La sortie est complémentée »	A  S	A  S	$S = \overline{A}$	<table><tr><th>A</th><th>S</th></tr><tr><td>0</td><td>1</td></tr><tr><td>1</td><td>0</td></tr></table>	A	S	0	1	1	0									
A	S																		
0	1																		
1	0																		
ET « La sortie vaut vrai si l'ensemble des entrées valent vrai »	A  B S	A  B S	$S = A \cdot B$	<table><tr><th>B</th><th>A</th><th>S</th></tr><tr><td>0</td><td>0</td><td>0</td></tr><tr><td>0</td><td>1</td><td>0</td></tr><tr><td>1</td><td>0</td><td>0</td></tr><tr><td>1</td><td>1</td><td>1</td></tr></table>	B	A	S	0	0	0	0	1	0	1	0	0	1	1	1
B	A	S																	
0	0	0																	
0	1	0																	
1	0	0																	
1	1	1																	
OU « La sortie vaut vrai si au moins une des entrées vaut vrai »	A  B S	A  B S	$S = A + B$	<table><tr><th>B</th><th>A</th><th>S</th></tr><tr><td>0</td><td>0</td><td>0</td></tr><tr><td>0</td><td>1</td><td>1</td></tr><tr><td>1</td><td>0</td><td>1</td></tr><tr><td>1</td><td>1</td><td>1</td></tr></table>	B	A	S	0	0	0	0	1	1	1	0	1	1	1	1
B	A	S																	
0	0	0																	
0	1	1																	
1	0	1																	
1	1	1																	
NON-ET (NAND) « La sortie vaut faux si l'ensemble des entrées valent vrai »	A  B S	A  B S	$S = \overline{A \cdot B}$	<table><tr><th>B</th><th>A</th><th>S</th></tr><tr><td>0</td><td>0</td><td>1</td></tr><tr><td>0</td><td>1</td><td>1</td></tr><tr><td>1</td><td>0</td><td>1</td></tr><tr><td>1</td><td>1</td><td>0</td></tr></table>	B	A	S	0	0	1	0	1	1	1	0	1	1	1	0
B	A	S																	
0	0	1																	
0	1	1																	
1	0	1																	
1	1	0																	
NON-OU (NOR) « La sortie vaut faux si au moins une des entrées vaut vrai »	A  B S	A  B S	$S = \overline{A + B}$	<table><tr><th>B</th><th>A</th><th>S</th></tr><tr><td>0</td><td>0</td><td>1</td></tr><tr><td>0</td><td>1</td><td>0</td></tr><tr><td>1</td><td>0</td><td>0</td></tr><tr><td>1</td><td>1</td><td>0</td></tr></table>	B	A	S	0	0	1	0	1	0	1	0	0	1	1	0
B	A	S																	
0	0	1																	
0	1	0																	
1	0	0																	
1	1	0																	
OU exclusif (XOR) « La sortie vaut vrai si les deux entrées sont différentes »	A  B S	A  B S	$S = A \oplus B$	<table><tr><th>B</th><th>A</th><th>S</th></tr><tr><td>0</td><td>0</td><td>0</td></tr><tr><td>0</td><td>1</td><td>1</td></tr><tr><td>1</td><td>0</td><td>1</td></tr><tr><td>1</td><td>1</td><td>0</td></tr></table>	B	A	S	0	0	0	0	1	1	1	0	1	1	1	0
B	A	S																	
0	0	0																	
0	1	1																	
1	0	1																	
1	1	0																	

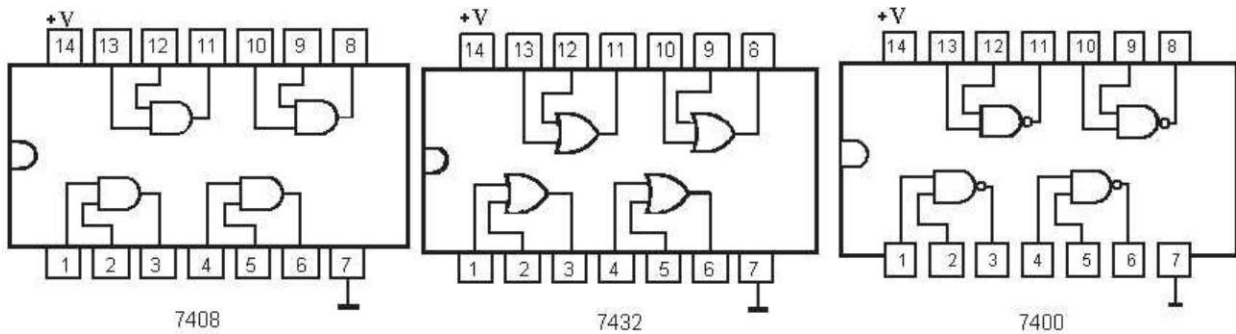
Remarque

Un porte logique peut avoir plus de deux entrées.

a) En pratique

Les portes logiques se trouvent communément sous la forme de circuits intégrés contenant plusieurs portes. Parmi les plus courants, nous retrouvons :

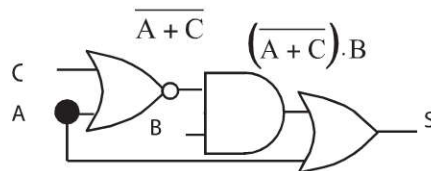
- Le 7408 est un circuit contenant quatre portes ET à 2 entrées ;
- Le 7400 est un circuit contenant quatre portes ET-NON à 2 entrées ;
- Le 7432 est un circuit contenant quatre portes OU à 2 entrées.



Soit la fonction logique suivante à réaliser en porte NAND.

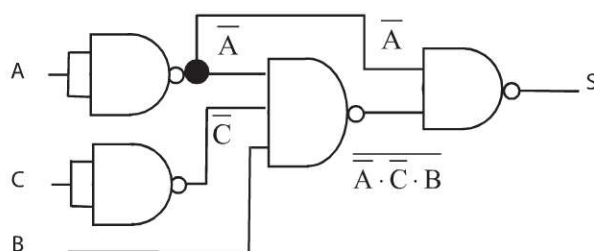
$$S = (\overline{A + C}) \cdot B + A$$

La réalisation directe conduit à :



L'opérateur NAND étant universel, il est possible de traduire n'importe quelle équation logique exclusivement à l'aide de ces portes en appliquant les théorèmes de De Morgan.

$$\begin{aligned} S &= (\overline{A + C}) \cdot B + A \\ &= \overline{\overline{A + C}} \cdot B + A \\ &= \overline{\overline{A} \cdot \overline{C}} \cdot B + A \\ &= \overline{\overline{\overline{\overline{A} \cdot \overline{C}} \cdot B}} + A \end{aligned}$$



La même démarche peut être adoptée pour une réalisation en portes NOR.

55 Bascules, logique séquentielle

Mots clés

Bascules RS, JK, D, latch, compteur, horloge.

1. EN QUELQUES MOTS

De même que les portes logiques, les bascules sont des composants numériques qui fonctionnent avec des signaux logiques à deux états (0 ou 1). La différence réside dans le fait que l'état de sortie d'une bascule dépend non seulement des entrées mais également de son état précédent : c'est ce que l'on appelle la logique séquentielle.

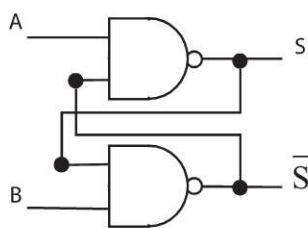
Les bascules sont à la base des systèmes tels que les mémoires ou les compteurs.

2. CE QU'IL FAUT RETENIR

a) Bascule RS

Cette bascule, la plus élémentaire, est constituée de deux portes NAND dont les sorties sont rebouclées en entrée tel que décrit sur la figure.

Posons $S = Q$. Soit Q_n l'état courant de la sortie et Q_{n+1} l'état futur. La table de vérité se réécrit de façon condensée :



A	B	S	$\bar{S}$
0	0	1	1
0	1	1	0
1	0	0	1
1	1	0	1
1	1	1	0

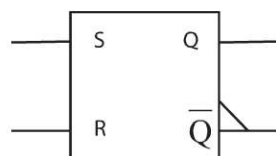
A	B	Q_{n+1}
0	X	1
0	1	0
1	0	Q_n

L'état X correspond à un niveau indéterminé (quelconque).

b) Bascule RS

Sur la base précédente, la bascule RS a un fonctionnement similaire :

- lorsque S seulement est à 1 (Set), la sortie est à 1 ;
- lorsque R seulement est à 1 (Reset), la sortie est à 0 ;
- si R et S sont à 0, l'état de la sortie est mémorisé ;
- si R et S sont à 1, l'état de la sortie est indéterminé.

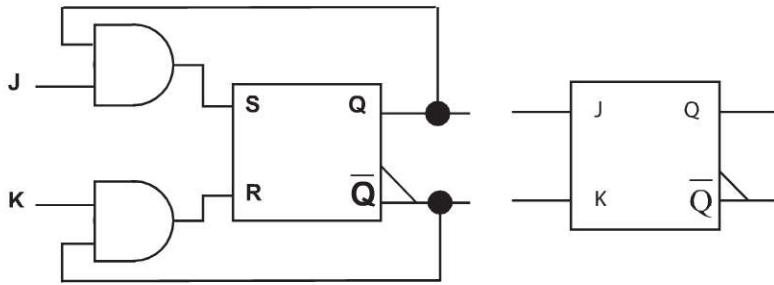


S	R	Q_{n+1}
0	0	Q_n
0	1	0
1	0	1
1	1	X

c) Bascule JK

La bascule JK est constituée d'une bascule RS et de deux portes ET.

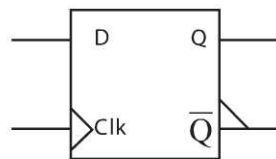
Elle permet de lever l'indétermination lorsque $R = S = 1$.



J	K	Q_{n+1}	Fonction
0	0	Q_n	Mémorisation
0	1	0	Mise à 0
1	0	1	Mise à 1
1	1	$\overline{Q_n}$	Inversion (toggle)

d) Bascule D

La bascule D (pour *Delay* - retard) est une bascule JK synchrone (avec entrée d'horloge) telle que $J = \overline{K} = D$. Elle stocke une valeur binaire tant qu'elle n'a pas reçu de coup d'horloge.



D	clk	Q_n
0	—	0
1	—	1
x	x	Q_{n-1}

Remarques

- Le symbole  signifie que la bascule est synchronisée sur front montant de l'horloge (passage de l'état 0 à l'état 1).
- Dans le cas contraire, l'entrée d'horloge  signifie que la bascule est synchronisée sur front descendant (passage de l'état 1 à l'état 0).

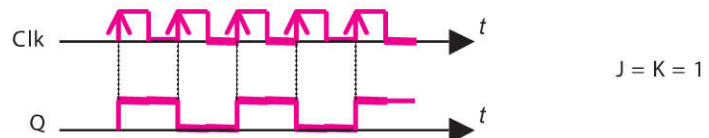
e) Latch

Une bascule D sans entrée d'horloge est un latch (verrou).

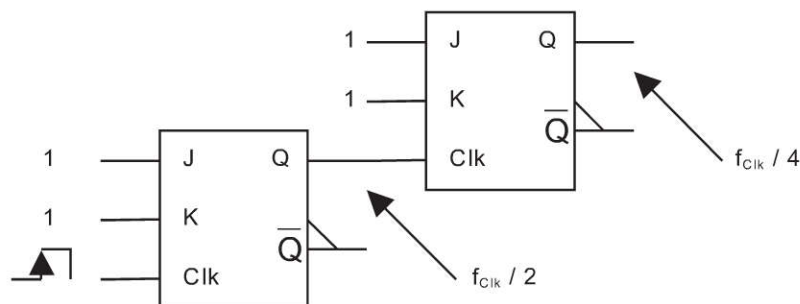
3. EN PRATIQUE

a) Compteur et division de fréquence

Soit une bascule JK synchrone (munie d'une entrée d'horloge) en mode toggle : ses deux entrées J et K sont à 1.



Nous voyons intuitivement apparaître deux applications à partir de ce principe : la division de fréquence d'horloge par 2^N (N étant le nombre de bascules) et les compteurs binaires basés sur le même schéma.



56 Circuits logiques programmables

Mots clés

ASIC, numérique, circuits logiciels et matériels, PAL, CPLD, FPGA, cycle de développement.

1. EN QUELQUES MOTS

Classiquement, pour réaliser un système numérique (horloge LCD, chenillard...), il faut définir les entrées et les sorties, décomposer le système et sélectionner les composants utiles à chaque sous-système. Ces composants sont ensuite interconnectés et le système est testé.

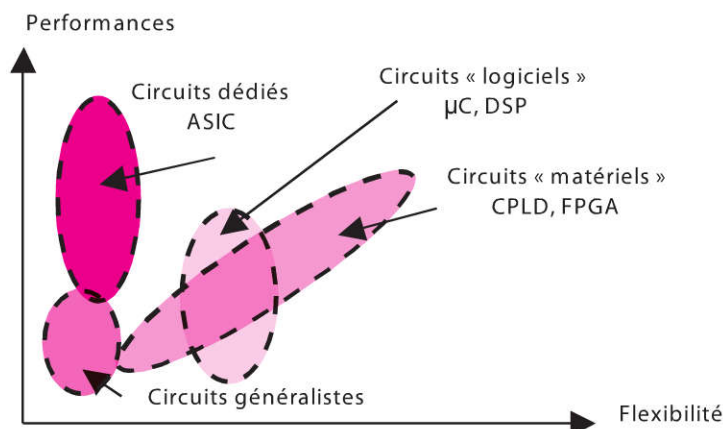
Cette méthode devient difficilement applicable dans un contexte industriel où les concepteurs de systèmes électroniques numériques doivent composer avec différentes contraintes : les temps et les coûts de développement doivent être réduits alors que les fonctions numériques à concevoir sont de plus en plus complexes. C'est ainsi que ces dernières années sont apparus de nouveaux composants et de nouvelles techniques de conception en électronique numérique.

2. CE QU'IL FAUT RETENIR

Les circuits intégrés numériques que l'on trouve sur les circuits imprimés peuvent être classés selon trois familles différentes.

- Les **circuits généralistes**, tels que les portes logiques, compteurs, bascules, mémoires, permettent de répondre à de nombreux besoins. Ces circuits, très répandus, sont bon marché mais ne sont pas forcément optimisés pour un système en particulier.
- Les **circuits spécifiques** ou **ASIC** (*Application Specific Integrated Circuits*) utilisent les technologies de pointe pour être parfaitement adaptés à un besoin particulier. En contrepartie, les phases de conception sont longues et la fabrication est coûteuse.
- Les **circuits logiques programmables (CLP)** proposent un compromis entre les circuits généralistes et les ASIC : ils sont très répandus et donc relativement bon marché tout en étant configurables. Ils peuvent donc répondre à de nombreux besoins. Parmi eux, distinguons deux sous-familles : les circuits logiciels (processeurs, microcontrôleurs [fiche 60](#)) et les circuits matériels (PAL, CPLD, FPGA...)([fiche 57](#)).

De façon très relative, nous pouvons situer ces familles les unes par rapport aux autres sur un graphique illustrant leur compromis entre performances (nombre d'opérations par seconde) et flexibilité (capacité à effectuer des fonctions différentes).



3. EN PRATIQUE

Les logiciels de conception de systèmes numériques pour CLP permettent de passer de la description de l'application au composant de façon optimisée, transparente, rapide et automatisée. Voici les points essentiels des cycles de développement des circuits logiciels et matériels.

a) Les circuits logiciels

► Modélisation du système numérique

L'utilisateur a le choix entre plusieurs langages source. Le langage assembleur, très proche du matériel, permet de décrire des opérations élémentaires (déplacement, calcul, modification du déroulement du programme, comparaison pour les principales). C'est un langage machine interprétable par l'humain. Pour plus de souplesse d'utilisation et de portabilité, des langages de plus haut niveau ont été intégrés aux outils de développement tels que le BASIC ou le C, ce dernier étant le plus répandu.

► Compilation

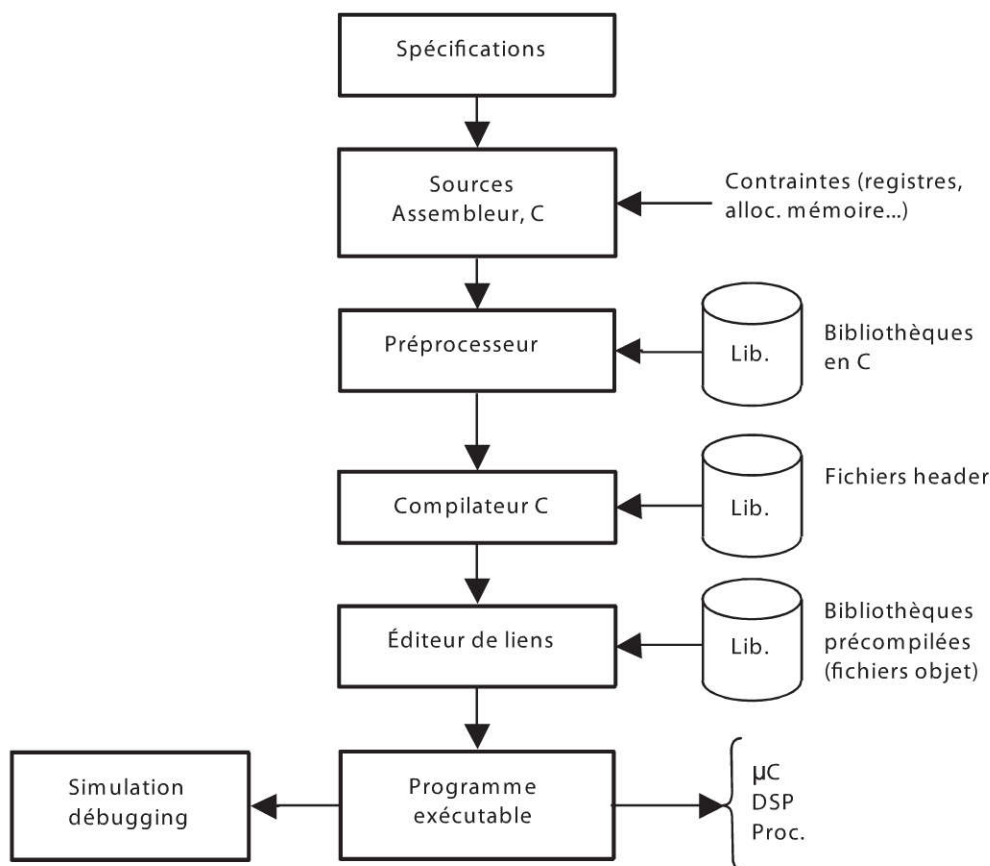
La compilation transforme le programme en langage source (assembleur, C...), en langage machine (langage cible), qui est une liste d'instructions binaires interprétables par le processeur.

► Linker

Le linker, ou éditeur de liens, fait le lien entre les différentes parties du code compilé et les espaces mémoire du CLP.

► Programmation

Le fichier binaire issu de la compilation doit être envoyé sur le circuit. Cela se fait soit en retirant le CLP du circuit, soit directement sur le circuit par l'intermédiaire d'un câble de programmation.



► Simulation, débbuging

La simulation consiste à faire « tourner » le programme sur ordinateur en tenant compte du fonctionnement du composant cible : cela permet d'observer l'évolution des différents registres et variables et ainsi de tracer les dysfonctionnements. Le débbuging se fait à l'aide d'émulateurs in-circuit qui permettent d'ajouter une interface de communication entre le composant sous test et un ordinateur.

b) Les circuits matériels

Les différentes étapes décrites ici partent du modèle du système : c'est la seule étape qui nécessite vraiment l'intervention de l'utilisateur.

► Modélisation du système numérique

L'utilisateur a le choix entre plusieurs méthodes, et donc différentes interfaces de saisie, pour décrire son application. Parmi celles-ci, la **description schématique** consiste à créer un schéma électronique du système : des composants virtuels sont placés et reliés ensemble.

Le codage en mode texte utilisant un **langage de description de matériel** ou HDL (*Hardware Description Language*) est une méthode plus abstraite : c'est avant tout le comportement du système qui est décrit. Les principaux HDL sont le VHDL (fiches 58 et 59), Verilog et Abel-HDL.

► Synthèse logique

La synthèse logique est l'étape traduisant en **blocs logiques élémentaires** interconnectés les fichiers du projet. L'utilisateur peut éventuellement agir sur la façon dont la synthèse s'effectue. Par exemple, il peut choisir d'optimiser le système en fréquence de fonctionnement ou en quantité de ressources logiques utilisées.

► Simulation fonctionnelle

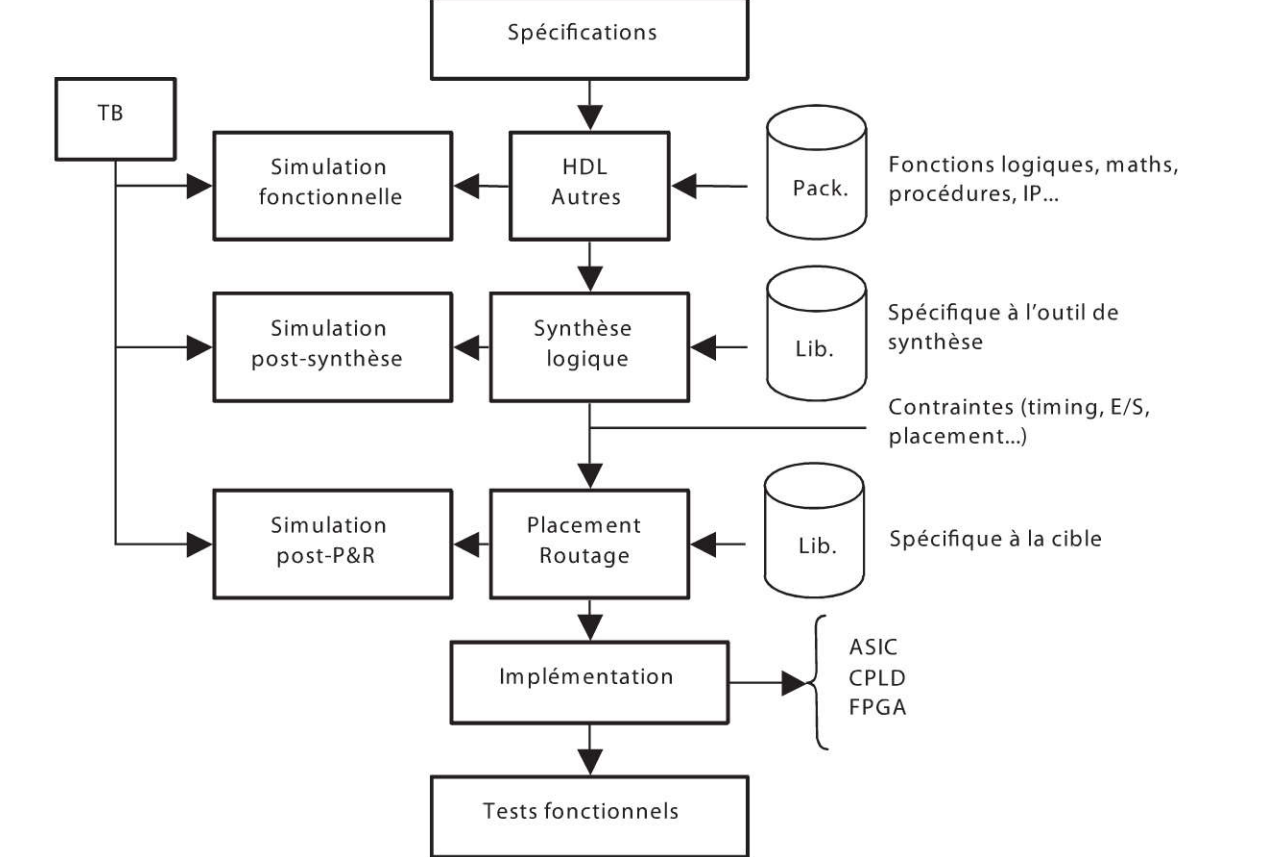
Il est important de pouvoir observer l'évolution des différents signaux d'un système soit pour le tester virtuellement avant de programmer le composant cible, soit pour diagnostiquer l'origine d'un dysfonctionnement. La simulation fonctionnelle est donc utile pour **tester le circuit** avant son implantation ou pour **détecter des erreurs de conception**. Elle ne tient pas compte des contraintes et aléas de fonctionnement liés au composant ciblé tel que la fréquence maximale d'horloge, les possibilités de routage, les ressources disponibles ou les temps de transfert réel dans les circuits logiques.

► Placement-routage (P&R)

Le placement-routage optimise l'implantation du système modélisé dans le composant cible. Au même titre qu'un concepteur de circuit imprimé devant prévoir l'implantation des composants électroniques et dessiner les pistes reliant ces composants, le placement-routage associe les pattes d'entrées/sorties aux signaux du modèle, réserve les différentes ressources logiques utiles (**placement**) et interconnecte ces mêmes ressources (**routage**).

L'effort d'optimisation (en surface ou en vitesse) peut être ajusté par l'utilisateur. La simulation post-routage produit un rapport avec les ressources utilisées et les temps de propagation de broche à broche.

La phase d'implémentation consiste à **configurer le composant**. Pour cela, un fichier de programmation est généré à la suite de la phase de placement-routage. Il contient toutes les données nécessaires à la configuration du composant : niveaux logiques des entrées/sorties, interconnexions entre les cellules logiques, fonctions combinatoires des LUT, utilisation de ressources câblées...



57 FPGA

Mots clés

CLB, bloc logiques configurables, ressources dédiées, JTAG, interconnexion.

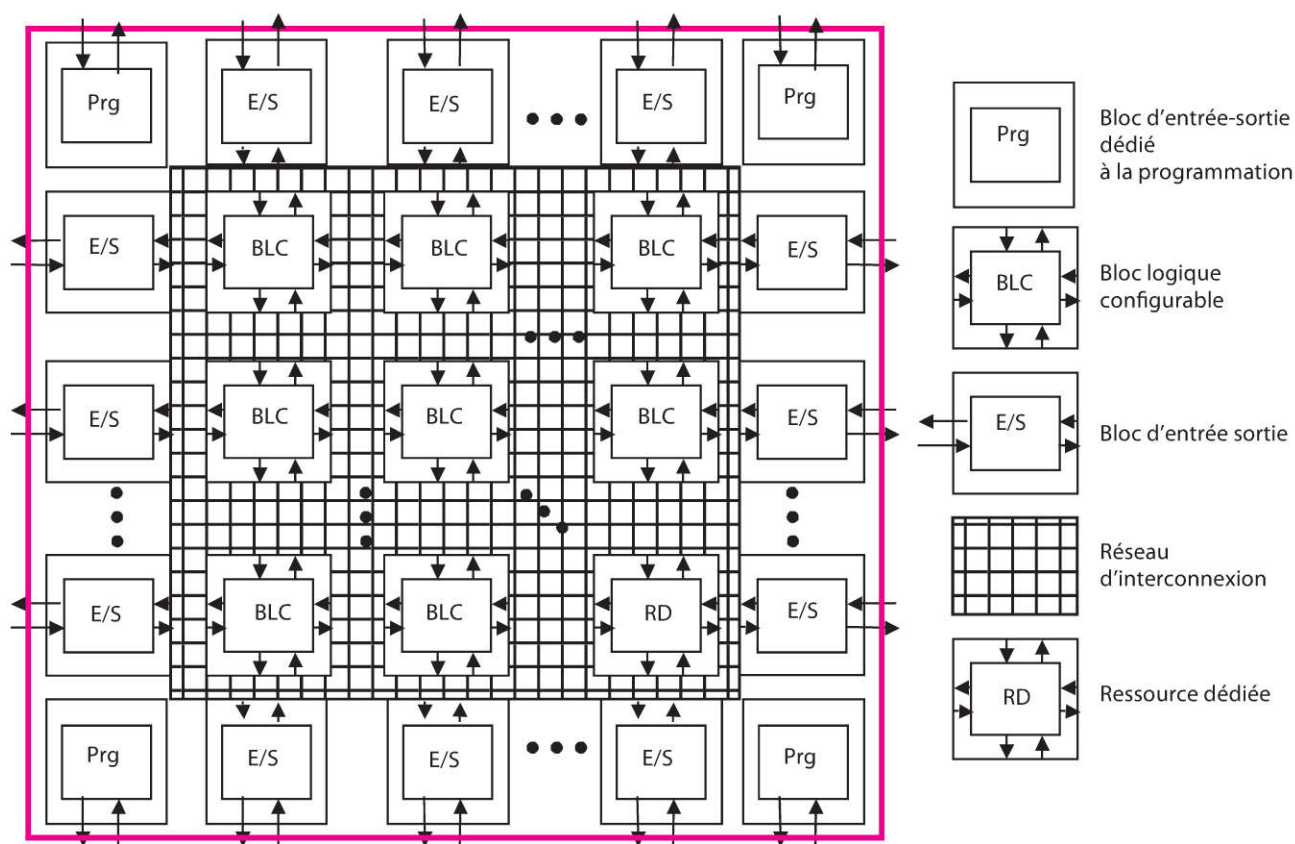
1. EN QUELQUES MOTS

Les FPGA (*Field Programmable Gate Arrays* ou réseaux de portes logiques programmables par l'utilisateur) font partie de la famille des circuits logiques programmables. Ce sont des circuits intégrés contenant un très grand nombre de portes logiques organisées en blocs logiques configurables et interconnectables. « Programmer » un FPGA, c'est configurer et interconnecter des blocs logiques.

2. CE QU'IL FAUT RETENIR

Les FPGA sont constitués des éléments suivants.

- Blocs d'entrées/sorties « IOB » configurables et accès de programmation du composant.
- Blocs logiques configurables.
- Lignes d'interconnexions.
- Ressources dédiées : multiplieurs, SRAM, etc...



a) Vue simplifiée

De façon simplifiée, il est possible de considérer un FPGA comme la superposition de trois couches.

- Les blocs logiques, les entrées/sorties et les ressources dédiées sont disposés sur la couche supérieur.
- Le réseau d'interconnexion permet d'effectuer les branchements entre les entrées/sorties, les blocs logiques et les ressources dédiées éventuellement utilisées.
- Une couche de mémoire configure les interconnexions, les entrées/sorties et les ressources logiques. Cette mémoire peut être volatile (elle s'efface lors de la mise hors tension du composant) ou bien antifusible (elle est programmée une fois pour toutes). Sa taille dépend directement de la « taille » du FPGA.

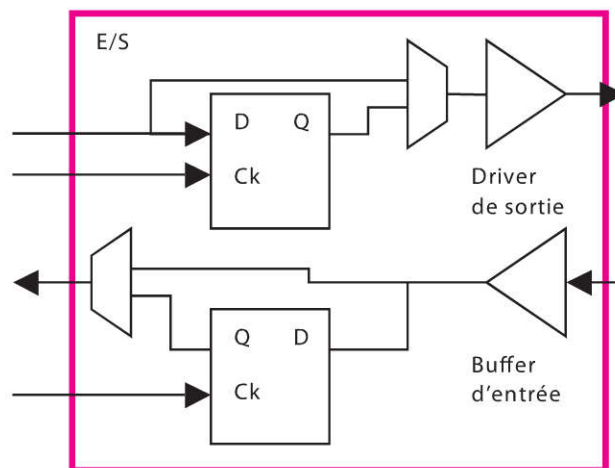
Remarque

Un antifusible, contrairement aux fusibles, permet d'établir une connexion lorsqu'un fort courant est appliqué.

b) Blocs d'entrée/sortie

Les blocs d'entrée/sortie peuvent être configurés, sauf exception, en tant qu'entrée, sortie ou accès bidirectionnel. La mémorisation d'une donnée entrante ou sortante peut se faire dans une bascule D : cela permet de synchroniser plusieurs signaux sur une horloge. L'utilisation ou non de ces points de mémorisation est assurée par des multiplexeurs.

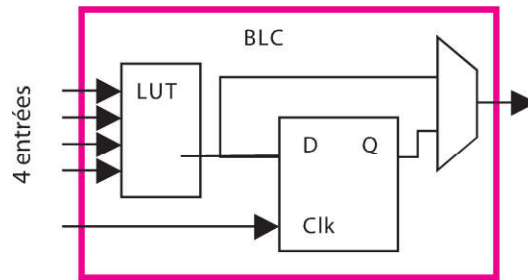
Les niveaux logiques sont configurables par l'utilisateur en fonction des tensions d'alimentation du FPGA. Ainsi, il est possible de configurer un accès en fonctionnement LVTTTL, LVCMOS, etc. et ce, généralement, avec des plages de tension allant de 0 V à 2,5 V ou 3,3 V pour les valeurs les plus typiques. Cette fonction d'adaptation des niveaux de tension est assurée par un driver (sortie) ou par un buffer (entrée).



c) Blocs logiques programmables

Le principe des blocs logiques programmables ou configurables (BLC ou CLB pour *Configurable Logic Block*) est au cœur du fonctionnement des FPGA (et par extension des CPLD). À partir de

plusieurs blocs génériques et du réseau d'interconnexion, il est possible d'obtenir toute fonction séquentielle ou combinatoire (arithmétique, mémoire, processeur...).



La partie combinatoire est constituée d'une table (LUT) comportant généralement quatre entrées et une sortie. Les quatre entrées de la LUT (*Look-Up Table*) permettent d'adresser 16 cellules mémoires de 1 bit chacune. Il est donc possible de générer grâce à elle n'importe quelle fonction combinatoire de un, deux, trois ou quatre variables binaires.

La partie séquentielle est constituée d'une bascule D pour la mémorisation. Le résultat présent en sortie de la LUT peut être envoyé sur cette bascule pour synchroniser les données au sein du système. Son utilisation dépend d'un multiplexeur qui sélectionne soit la sortie de la bascule, soit directement la sortie de la LUT.

d) Gestion d'horloge

Des entrées spécialisées sont prévues pour recevoir les signaux d'horloge et les distribuer à l'intérieur du circuit. Ces lignes d'interconnexions sont distinctes des autres afin de limiter les perturbations.

Les DCM (*Digital Clock Managers*), dont le principe est basé sur celui des PLL (fiche 64), sont des blocs internes pour la gestion du signal d'horloge. Ils peuvent générer des déphasages et des fréquences d'horloge multiples de la fréquence du signal d'horloge externe, tout en éliminant le « jitter ». Le jitter, ou gigue de phase, représente un décalage des fronts d'un signal binaire. Ce phénomène a une influence non négligeable, notamment pour le pilotage de convertisseurs A/N ou N/A (fiches 62 et 63).

La multiplication de fréquence d'horloge est limitée à un certain domaine de fréquences. Le coefficient multiplicateur est un rapport rationnel dont le choix des opérandes est également restreint.

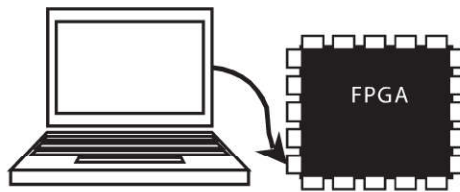
e) Ressources dédiées

Nous avons vu que les blocs logiques configurables sont capables d'assurer toute fonction numérique. Cependant, certaines fonctions, parfois très utilisées telles que la multiplication ou la mémoire, mobilisent une grande quantité de ressources. Pour cette raison, des blocs de fonctions dédiés sont presque toujours mis à la disposition du concepteur pour optimiser le système en termes de ressources et de vitesse de fonctionnement. Les fonctions proposées dépendent de la marque, de la famille et de la « taille » du FPGA. Parmi celles-ci, on peut retrouver de la mémoire, des multiplieurs ou des multiplieurs-accumulateurs (indispensables pour la FFT et la plupart des filtres numériques), des blocs de gestion d'horloge (DCM), des microcontrôleurs ou des processeurs, des blocs de gestion de port série...

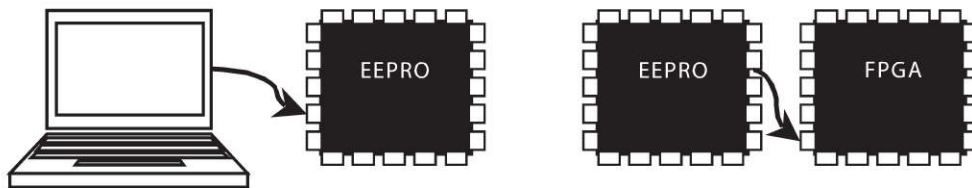
f) Programmation

Afin de programmer un FPGA, il faut lui envoyer une chaîne de bits de configuration. Cette chaîne spécifie l'état des lignes d'interconnexion, des cellules logiques, des ressources dédiées et des entrées/sorties. Le FPGA est programmé soit de façon temporaire si la mémoire interne est volatile, soit de façon permanente s'il utilise une technologie EPROM, fusible ou antifusible. Cette catégorie de FPGA n'est pas reprogrammable, ce qui peut trouver un intérêt dans des

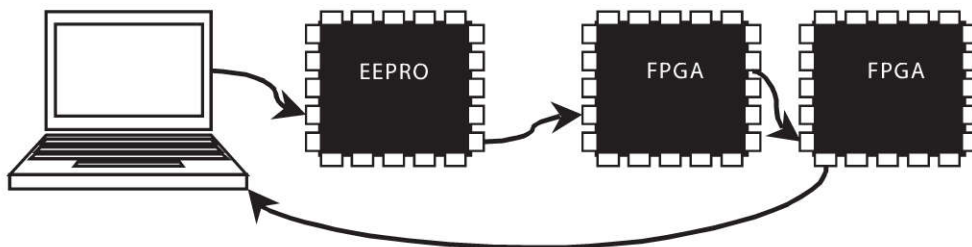
applications où le fonctionnement ne doit pas être modifié, où les perturbations électromagnétiques ou les radiations sont fortes et pour certaines productions en grande série.



Pour les FPGA reprogrammables, afin de conserver une ou plusieurs configurations du composant à la mise hors tension du système, il convient de stocker les données de configuration dans une mémoire non volatile de type EEPROM (*Electrically Erasable Programmable Read-Only Memory* – Mémoire morte programmable et effaçable électriquement) ou Flash. Cela évite de recourir à l'utilisation d'un ordinateur après chaque coupure d'alimentation.



La programmation directe d'un FPGA et de certaines mémoires de configuration externes se fait par connexion JTAG (*Joint Test Action Group*), également appelée Boundary Scan (« scrutation des frontières »). Initialement dédié au test des circuits numériques, ce protocole permet de programmer des composants (FPGA, microcontrôleurs). Il est possible de chaîner plusieurs composants programmables en JTAG : une seule connexion suffit à gérer plusieurs circuits intégrés. Le transfert est rythmé par une horloge.



3. EN PRATIQUE

Les performances d'un FPGA s'évaluent en fonctions de plusieurs paramètres. Le premier est le nombre de portes logiques ou le nombre de blocs logiques configurables qui est un indicateur des ressources mises à disposition. L'inventaire des ressources dédiées qui complètent les CLB est essentiel : a-t-on besoin d'un ou plusieurs processeurs ? de beaucoup de mémoire ? d'effectuer des opérations arithmétiques lourdes ?

La vitesse à laquelle les opérations peuvent être cadencées va beaucoup dépendre du placement routage, mais elle repose aussi sur la technologie du FPGA et son *speed grade* (catégorie de rapidité de fonctionnement).

Les autres critères de choix vont concerner la puissance consommée, la taille du boîtier et le nombre d'entrées-sorties disponibles.

58 Le langage VHDL

Mots clés

Entity, architecture, std_logic, process, begin, end, description matérielle.

1. EN QUELQUES MOTS

Le VHDL (*Very high-speed integrated circuit Hardware Description Language* ou langage de description matérielle pour circuits intégrés très rapides) n'est pas à proprement parler un langage de programmation, mais il est un langage de description de matériel électronique. Il a été initié par le DoD (*Department of Defense*, ou Défense américaine) en 1980 puis il est devenu standard IEEE en 1987 et revu en 1993 (std IEEE 1076-1993). En 1999, il a été élargi aux signaux analogiques (VHDL-AMS : std IEEE 1076.1). Le langage VHDL est un des points d'entrée de la conception de systèmes à base de FPGA. Il permet de modéliser des applications au même titre que les langages Verilog ou ABEL, moins répandus, et qui ne sont pas présentés ici. Cette conception peut se faire sous la forme de schémas logiques : cette approche, intuitive, est facile d'accès et bénéficie de l'utilisation de blocs de base disponibles dans une bibliothèque... Le manque de généricité et de lisibilité en font une approche peu adaptée à une conception hiérarchique et à la conception de systèmes un peu complexes : machines d'état, communication, arithmétique... La conception utilisant des langages de description de matériel permet de décrire un système de façon simple, structurée, portable et réutilisable.

2. CE QU'IL FAUT RETENIR

a) Fichiers

En VHDL, tout est inclus dans des bibliothèques : l'unité la plus englobante est la bibliothèque (**library**). Un composant VHDL est constitué de son entité (**entity**) représentant son interface avec l'extérieur, et de son **architecture** représentant la description interne du composant. Des fonctions et des procédures ou d'autres déclarations peuvent être ainsi regroupées sous la forme d'un paquetage (**package**). La **configuration** des instances (occurrences) des composants peut se faire soit directement à l'instanciation des composants, soit sous la forme d'une unité de **configuration**.

b) Types

Les types prédéfinis scalaires sont :

- integer codé sur 32 bits (parfois 64 bits) ;
- natural codé sur 31 bits ;
- boolean (false, true) ;
- bit ('0', '1') ;
- char (NUL, ...'a', 'b', 'c'...) ;
- time $-(231-1)$ à $-231-1$ (integer ou real) ;
- real (flottant dépendant de l'implémentation).

Les types prédéfinis complexes sont :

- bit_vector (tableau de bit) ;
- string (tableau de caractères char).

Les types `std_ulogic` et `std_logic` ne sont pas définis par le standard VHDL à proprement parler mais ils sont très utilisés dans la conception de systèmes. Pour pouvoir les manipuler, il faut faire appel à la bibliothèque `IEEE_1164` en en-tête du fichier VHDL de la façon suivante :

```
library IEEE;
use IEEE.std_logic_1164.all;
```

Cette syntaxe est à utiliser pour tout appel à la bibliothèque. Les types `std_ulogic` et `std_logic` sont une spécificité de la description de systèmes numériques. Ils peuvent prendre les neuf valeurs suivantes :

'U'	Uninitialized	Non initialisé
'X'	Forcing Unknown	Conflit fort
'0'	Forcing 0	0 fort
'1'	Forcing 1	1 fort
'Z'	High Impedance	Haute impédance
'W'	Weak Unknown	Conflit faible
'L'	Weak 0	0 faible
'H'	Weak 1	1 faible
'.'	Don't care	Sans influence

Un type peut être restreint à un sous-domaine. Le mot clé **range** définit un sous-intervalle de valeurs ordonné. Par exemple, le sous-type `chiffre` défini ici est une restriction des valeurs de 0 à 9 du type `integer` :

```
subtype chiffre is integer range 0 to 9 ; -- ordre croissant
```

L'énumération permet de créer un type en listant de manière ordonnée les identificateurs ou caractères du type. Dans cet exemple, `couleur` ne peut prendre que les valeurs « rouge », « vert » et « bleu » :

```
type couleur is (rouge, vert, bleu)
```

Les tableaux sont des listes ordonnées de valeurs d'un même type. Ici, le type `matrice` est un tableau de 10 par 20 valeurs booléennes :

```
type matrice is array (0 to 9, 0 to 19) of boolean ;
```

Le type `bit_vector`, déclaré comme un tableau de valeurs binaires, permet de définir un vecteur de bits (c'est-à-dire potentiellement une case mémoire ou un bus)

```
type bit_vector is array (natural range <>) of bit ;
```

Le sous-type `octet` est défini à partir de `bit_vector` et restreint le nombre de bits à 8:

```
subtype octet is bit_vector(7 downto 0) ;
```

c) Opérateurs

Les opérateurs permettent d'effectuer des opérations logiques (fiches 52 et 54), arithmétiques ou relationnelles. Dans le cas général, ces opérations s'appliquent à deux opérandes, sauf pour NOT qui inverse ou complémente un opérateur logique et '-' qui peut représenter l'inverse d'un nombre aussi bien que la soustraction.

Logiques		Relationnels		Arithmétiques	
AND	et	=	égal	+	addition
OR	ou	/=	différent	-	soustraction
NAND	non et	<	inférieur	*	multiplication
NOR	non ou	<=	inférieur ou égal	/	division
XOR	ou exclusif	>	supérieur	**	puissance
NOT	inversion	>=	supérieur ou égal	MOD	reste de la division de deux entiers, le signe dépend de l'opérateur
				REM	
				&	Concaténations

Remarque

La concaténation permet « d'agglomérer » deux éléments de même type. Par exemple, "bon" & "jour" donne "bonjour".

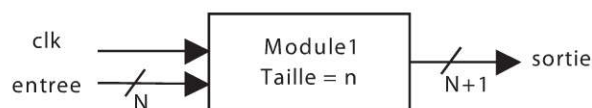
Les opérateurs logiques peuvent s'effectuer sur des éléments de type boolean, bit, std_ulogic et leurs dérivés. Les opérateurs relationnels renvoient une valeur booléenne. Les opérateurs arithmétiques s'appliquent à des types différents en fonction des bibliothèques appelées. Les plus courantes sont :

- library IEEE;
- use IEEE.std_logic_1164.all;
- use IEEE.numeric_std.all;
- use IEEE.std_logic_unsigned.all;
- use IEEE.std_logic_signed.all;
- use IEEE.std_logic_arith.all;
- use IEEE.math_real.all;

d) Entité et architecture

Une **entité** est une boîte noire indiquant le nom, les paramètres génériques (generic) et les entrées/sorties d'un module (port). Les paramètres génériques permettent de déclarer des constantes globales telles que la taille d'un bus par exemple. Les entrées/sorties ont un nom, un type et une direction qui peut être out (sortie), in (entrée), inout (bidirectionnel) ou buffer.

L'entité Module1 représenté sur le schéma suivant correspond au code VHDL ci-après en considérant que entree est une entrée sur 8 bits, et sortie, une sortie sur 9 bits conformément à l'initialisation du paramètre générique N.



```

ENTITY module1 IS
  GENERIC (
    N : integer := 8
  )
  PORT (
    entree : IN bit_vector(N-1 DOWNTO 0);
    Clk : IN bit;
    Sortie : OUT bit(N DOWNTO 0)
  );
END;
```

L'**architecture** contient ce qui se passe à l'intérieur de la boîte noire entité. Il existe deux façons de décrire un système. Une **vue structurelle** fait appel à des modules existants et les relie entre

eux. Une **vue comportementale** consiste à utiliser le langage pour indiquer les opérations que le système doit effectuer. Ces opérations peuvent se situer soit sur un **mode concourant**, tout ce passe en simultanée au même titre que les circuits intégrés placés sur un circuit imprimé fonctionnent tous en même temps, soit sur un **mode séquentiel**, les opérations s'effectuant les unes à la suite des autres, comme dans un processeur (fiche 60). Les opérations séquentielles sont codées à l'intérieur d'une zone de programme appelée **process**. Un process est un groupe délimité d'instructions, doté de trois caractéristiques essentielles : le process s'exécute à chaque changement d'état d'un des signaux auxquels il est déclaré sensible, les instructions du processus sont testées dans l'ordre où elles sont écrites, les modifications apportées aux valeurs de signaux par les instructions prennent effet à la fin du processus.

Les boucles et branchements, à l'image de celles utilisées dans les langages informatique type C, Pascal... ne sont utilisables que dans les process.

<pre> case (<2-bit select>) is when "00" => <instruction>; when "01" => <instruction>; when "10" => <instruction>; when "11" => <instruction>; when others => <instruction>; End case; </pre>	<pre> if <condition> then <instruction> elsif <condition> then <instruction> else <instruction> end if; with <choice_expression> select <nom> <= <expression> when <ch>, <expression> when <ch>, <expression> when others; </pre>	<pre> for <nom> in <limite_basse> to <limite_haute> loop <instruction>; <instruction>; end loop; while <condition> loop <instruction>; <instruction>; end loop; </pre>
--	---	---

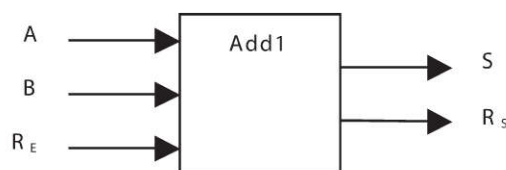
Remarques importantes

- Le VHDL n'est pas sensible à la casse (distinction minuscules/majuscules).
- L'indication '--' marque le début d'une ligne de commentaires.

3. EN PRATIQUE

Soit un additionneur dit « 1 bit » à coder en VHDL. Les équations logiques donnant la somme et la retenue, issues des entrées A, B et C, sont indiquées avec le schéma du module.

Boîte noire



Vue comportementale

$$S = A \oplus B \oplus C$$

$$R_S = A \cdot B \cdot \overline{R_E} + (A + B) \cdot R_E$$

ENTITY add1 IS

PORT(

```

  A : IN std_logic;
  B : IN std_logic;
  Re : IN std_logic;
  Rs : OUT std_logic;
  S : OUT std_logic;

```

END;

ARCHITECTURE gate OF add1a IS

SIGNAL D : std_logic;

BEGIN

```

  D <= A XOR B;
  S <= D XOR Re;
  Rs <= (A AND B AND NOT Re)
        OR ((A OR B) AND Re);

```

END gate;

Il est à noter que l'utilisation du signal D n'apporte rien au modèle de l'additionneur. Sa présence sert à illustrer l'utilisation d'un signal interne et qui, par conséquent, ne provient pas de l'extérieur du module.

59 Pratique du VHDL

Mots clés

Process, instanciation, liste de sensibilité, banc de test, instructions concourante, instruction séquentielle.

1. EN QUELQUES MOTS

Le langage VHDL peut être étudié comme un langage à part entière, sans se préoccuper de son utilisation finale. Dans le cadre de l'électronique, le VHDL sert à décrire le fonctionnement d'un système avec pour objectif de construire un circuit intégré numérique de type **ASIC** ou de **programmer un circuit logique programmable** tel qu'un **FPGA** (fiche 57). C'est cette deuxième finalité à laquelle un étudiant en électronique s'intéresse dans la plupart des cas.

Afin de détecter les erreurs de conception, il convient de savoir **simuler** l'application développée : cette étape permet simplement de visualiser l'évolution des valeurs internes à l'application sans avoir recours à des tests fastidieux sur le composant numérique lui-même.

2. CE QU'IL FAUT RETENIR

a) Utilisation des process

Dans un système électronique numérique, **tout fonctionne en simultané**. Par exemple, une porte logique n'attend pas qu'une bascule ailleurs dans le montage reçoive un coup d'horloge pour donner son résultat. Les opérations n'ont pas d'ordre au même titre que dans un processeur : en dehors d'un process et à l'intérieur d'une architecture, c'est-à-dire dans **le domaine concourant**, les lignes de code VHDL sont interchangeables.

Lorsqu'on a recours à une mémorisation de données (bascule), des instructions à exécuter de façon ordonnée et **séquentielle** telles que dans des machines d'état, les opérations effectuées dépendent d'une horloge, il faut recourir aux **process**. Un process est une portion de code VHDL dans laquelle les instructions sont séquentielles.

Remarque

Un process est lui-même vu comme une instruction concurrente, il peut donc être placé n'importe où dans l'architecture...

Dans cet exemple, clock constitue le signal présent dans la **liste de sensibilité** du process. C'est sur le signal clock que le process peut être déclenché, c'est-à-dire que les instructions sont exécutées lorsqu'un changement d'état survient sur l'un des signaux de la liste de sensibilité, en l'occurrence, clock. « ex » est le nom du process. La **variable** var est une valeur entière qui ne peut être vue que de l'intérieur du process. Contrairement aux signaux dont la valeur est actualisée en sortie d'un process, une variable est actualisée instantanément et s'affecte avec le symbole « := ».

Les instructions séquentielles sont situées entre les mots clés begin et end.

Le process présenté ici utilise une variable nommée var, de type std_logic_vector, déclarée avant le mot begin. Nous supposons que toto est un signal de même type et de même taille (tableau de 8 std_logic) et donc déclaré en tant que sortie de l'entity ou comme signal interne à l'architecture. Le signal reset, supposé de type std_logic, permet d'affecter la valeur prise par var à toto et de réinitialiser var. Si ces deux lignes étaient inversées, toto qui prend la valeur de var prendrait la valeur « 00000000 » : l'ordre des lignes de code a donc une influence à l'intérieur d'un process. Le

signal reset est synchrone à l'horloge : il faut avoir détecté un front montant sur clock pour tester ensuite la valeur de reset. Pour rendre asynchrone la remise à zéro, il faudrait ajouter le signal reset dans la liste de sensibilité du process ex et placer le test de la valeur de reset avant de tester le front de clock.

Ce process incrémente donc la valeur de var à chaque front montant de clock et ce, jusqu'à ce que reset soit actif (niveau '1' validé par clock). La valeur de var est alors copiée dans toto avant que celle-ci ne soit réinitialisée.

ex : process (clock)

VARIABLE var : std_logic_vector(7 DOWNT0 0);

begin

```

    if rising_edge(clock) then
        if reset = '1' then
            toto <= not var;
            var := "00000000";
        else
            var := var + '1';
        end if;
    end if;

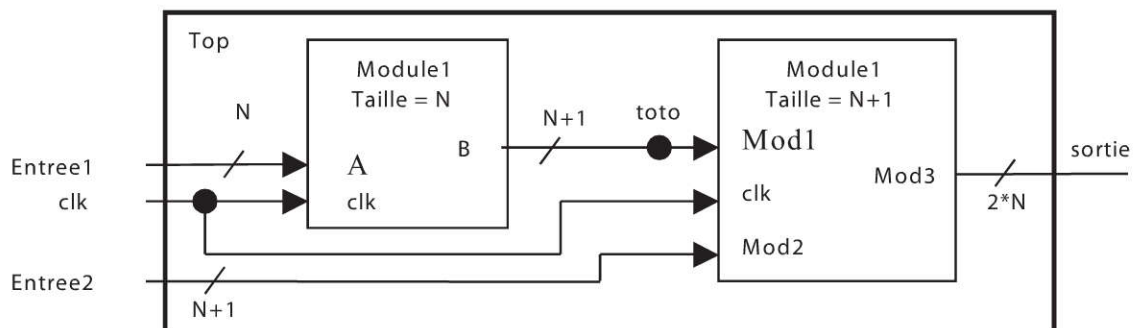
```

end process ex;

► Instanciation

Nous avons vu dans la [fiche 58](#) qu'il existe deux méthodes complémentaires pour décrire une application en VHDL : la vue comportementale (telle que l'exemple précédent) et la vue structurale. La vue structurale consiste à assembler différents modules entre eux. Le fait de réutiliser un module dans une architecture s'appelle l'**instanciation**. Il faut pour cela :

- avoir défini un module à instancier, c'est-à-dire au minimum son entity ;
- déclarer le module à utiliser en utilisant le mot clé **component** ;
- l'instancier à l'intérieur de l'architecture, ce qui consiste à reprendre le nom du module, fixer les paramètres génériques à l'aide des mots clés **generic map** (s'il y en a et si l'utilisateur veut utiliser d'autres paramètres que ceux déclarés par défaut) et associer les ports à des signaux ou à des ports de l'architecture utilisant le module en question avec les mots clés **port map**. L'association d'un paramètre générique avec une valeur ou d'un port avec un autre se fait à l'aide du symbole =>.



```

ENTITY top IS
    GENERIC(N : natural := 8)
    PORT( Entree1 : IN std_logic_vector(N-1 DOWNT0 0);
          Entree2 : IN std_logic_vector(N DOWNT0 0);
          Clk : IN bit;

```



```

        Sortie : OUT std_logic_vector(2*N DOWNT0 0));
END;
ARCHITECTURE struct OF top IS

    SIGNAL toto : std_logic_vector(N DOWNT0 0);
    COMPONENT module1
    GENERIC(taille : natural)
    PORT( A      : IN std_logic_vector(N-1 DOWNT0 0);
          Clk    : IN std_logic ;
          B      : OUT std_logic_vector(N DOWNT0 0));
    COMPONENT module2
    GENERIC(taille : natural)
    PORT( Mod1   : IN std_logic_vector(N-1 DOWNT0 0);
          Mod2   : IN std_logic_vector(N-1 DOWNT0 0);
          Clk    : IN std_logic;
          Mod3   : OUT std_logic_vector(N DOWNT0 0));

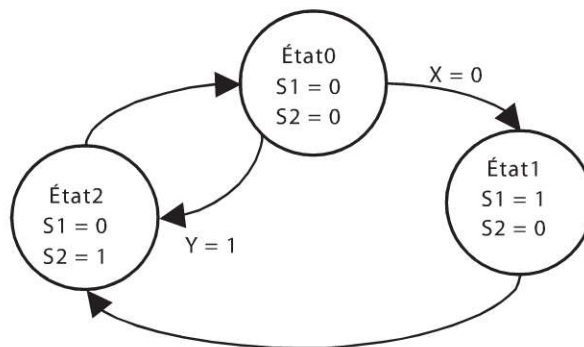
    BEGIN
        Inst1 : module1
        GENERIC MAP (taille => N)
        PORT MAP ( A=>Entree1, Clk=>Clk, B=>toto);
        Inst2 : module2
        GENERIC MAP (taille => N+1)
        PORT MAP ( Mod1=>toto, Clk=>Clk, Mod2=>entree2, Mod3=>sortie);
    END ARCHITECTURE struct ;

```

3. EN PRATIQUE

a) Machines d'état

Les machines d'état sont des méthodes de gestion de processus séquentiels. Elles sont très utilisées pour contrôler des protocoles de communication, des automates ou des algorithmes de calcul.



Dans l'exemple proposé ici et partiellement codé en VHDL, la machine est composée de trois états (État0 à État2, État0 est l'état initial), de deux sorties (S1 et S2) et de deux entrées X et Y. Une remise à zéro permet l'initialisation. Les indications Y = 1 et X = 0 sont des conditions de transition. Lorsqu'il n'y en a pas, la transition est inconditionnelle et l'entrée d'horloge seule fait avancer le processus. Pour résumer, une transition inconditionnelle se fait sur un coup d'horloge et une transition conditionnelle se fait si la condition est remplie.

```

ARCHITECTURE diag OF machine1 IS
    TYPE etat_3 IS (Etat0, Etat1, Etat2);
    SIGNAL etat,etat_suiv :etat_3 := Etat0;

BEGIN

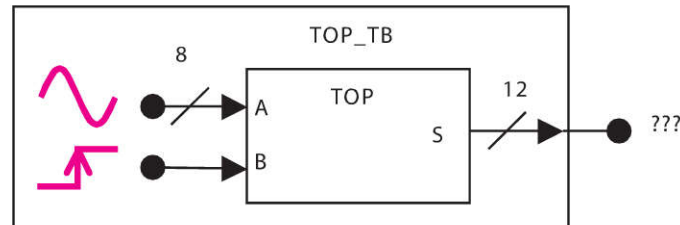
Synchro : PROCESS(reset, clk)
BEGIN
    If reset = '1' THEN
        etat <= Etat0;
    ELSIF rising_edge(clk) THEN
        etat <= etat_suiv ;
    END IF;
END PROCESS Synchro;

mef : process (etat, x, y)
BEGIN
    CASE etat IS
        WHEN Etat0 =>
            S1 <= '0';
            S2<= '0';IF Y='1' THEN
                etat_suiv <= Etat3;
            ELSIF X='0' THEN
                etat_suiv <= Etat2;
            ELSE
                etat_suiv <= Etat1;
            END IF;
        WHEN Etat1 =>
            S1 <= '1';
            S2 <= '0';
            etat_suiv <= Etat2;
        WHEN Etat2 =>
            S1 <= '1';
            S2<= '1';
            etat_suiv <= Etat3;
        WHEN others =>
            S1 <= '0';
            S2<= '0';
            etat_suiv <= Etat0;
    END CASE;
END process mef ;

```

b) Les bancs de test

Il est parfois indispensable de tester un module afin d'y apporter des corrections. À cette fin, le VHDL permet de décrire des **bancs de test** (ou *testbenches*, voire TB). Un banc de test n'a ni entrées, ni sorties. Par conséquent, l'**entité est vide**. Il instancie le ou les modules à tester et génère les signaux de test. Ces signaux sont soit décrits de façon comportementale, soit de façon structurelle en instanciant des modules de génération de signal, soit en lisant des fichiers de points grâce à la bibliothèque **stdio**. Voici un exemple dans lequel le banc de test `top_tb.vhd` teste le module `top`.



Il faut générer les signaux afin d'exciter les entrées du dispositif sous test. Parmi ceux-ci, les horloges sont très souvent indispensables pour cadencer les opérations. Voici un exemple de code VHDL pour générer un signal d'horloge `Clk` dont la période, nommée `period`, est du type `time` et prend ici pour valeur 10 ns, ce qui correspond à une fréquence d'horloge de 100 MHz. L'autre exemple montre comment générer un signal en dents de scie : la valeur de `ramp` s'incrémente toutes les 10 ns. Lorsque la valeur maximum est atteinte (N bits à '1'), `ramp` repart de zéro.

```
signal Clk : std_logic := '0';
constant period: time := 10ns;
BEGIN

pclk : process
begin
    Clk <= not(Clk);
    wait for period/2;
end process pclk;
[...]
END ARCHITECTURE;
```

```
signal ramp: std_logic_vector(N-1 downto 0)
:= (others => '0');
Constant period : time := 10ns;

BEGIN
pramp : process
begin
    ramp <= ramp + '1';
    wait for sample;
end process;
[...]
END ARCHITECTURE;
```

60 Processeurs

Mots clés

Processeur, instruction, plan mémoire, bus, contrôle, donnée, microcontrôleurs.

1. EN QUELQUES MOTS

Un système informatique est une carte électronique regroupant un ensemble de fonctionnalités pilotées par un microprocesseur. Il en existe deux principales architectures constituant un système informatique : l'architecture Von Neumann et l'architecture Harvard.

2. CE QU'IL FAUT RETENIR

a) L'architecture Von Neumann

► Système minimum

Le microprocesseur a besoin de ressources (composants) pour exécuter un programme. Les ressources minimales nécessaires sont :

- de la mémoire contenant le programme à exécuter (boîtiers mémoire de type ROM) ;
- de la mémoire vive (boîtiers mémoire de type RAM) pour stocker des résultats de ses opérations ;
- une interface pour communiquer (boîtiers de communication) avec l'utilisateur ou l'environnement.

Pour communiquer avec ses ressources, le microprocesseur est pourvu de bus nécessaires pour son dialogue avec les composants :

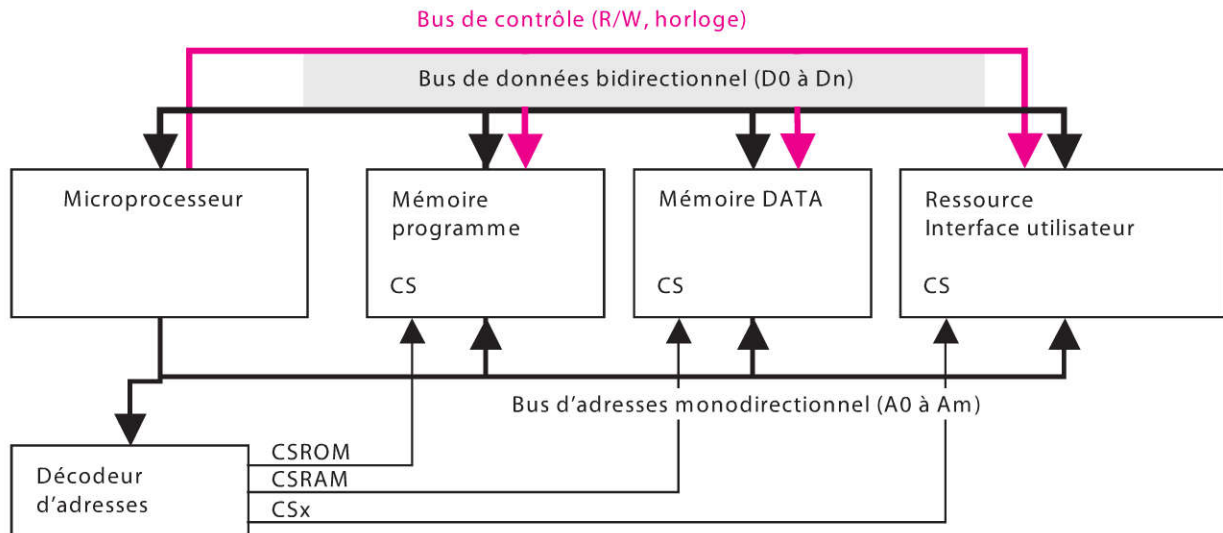
- un bus d'adresses permettant de choisir la case mémoire sur laquelle il va travailler ;
- un bus de données pour coder la valeur numérique de l'information échangée ;
- un bus de contrôle codant le sens de l'échange (lecture ou écriture) et contenant des signaux d'horloge.

L'activation d'un boîtier est confiée à un décodeur d'adresses chargé de l'activer selon son adresse dans le plan mémoire global du microprocesseur. Cette activation / désactivation se fait par un signal spécifique généralement actif à l'état bas (CS pour *chip select* ou CE pour *chip enable*) disponible sur chacun des boîtiers qui vient forcer son bus de données en état haute impédance (non connecté). Le décodeur d'adresse doit fournir autant de signaux CS (ou CE) qu'il y a de composants connectés sur son bus de données.

Cette représentation est universelle pour tout système informatique basé sur l'architecture Von Neumann. Seuls changent le nombre et la nature des composants connectés et par là même la structure interne et externe du décodeur d'adresses. Des limitations de performances sont présentes dans cette architecture. La présence d'un seul plan mémoire impose au microprocesseur de faire de nombreux accès dans ce plan pour :

- récupérer chaque instruction programme dans la mémoire ROM : une instruction est généralement codée sur plusieurs octets (jeux d'instruction CISC – *Complex Instruction Set Computer*) ;
- décoder l'instruction puis écrire (si nécessaire) les résultats en mémoire vive.

Ce qui prend de nombreux cycles (temps d'horloge) pour la récupération de l'instruction, son décodage et son exécution.



b) Le plan Mémoire

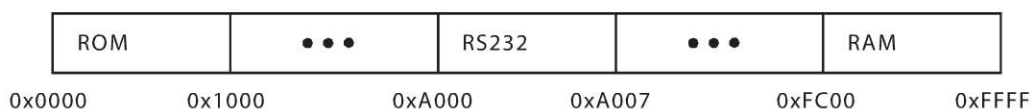
Il s'agit d'une représentation simplifiée de l'intégralité d'un plan électronique d'un système informatique pour son exploitation (programmation) ou pour sa conception. Il décrit :

- l'intervalle mémoire total adressable par le microprocesseur qui est fonction de la taille de son bus d'adresse : avec n fils, on peut adresser un intervalle de $[0 \text{ à } 2^{n-1}]$
- l'intervalle des adresses activant un composant
- la nature des composants implantés dans le SI

Si on se réfère au schéma précédant composé :

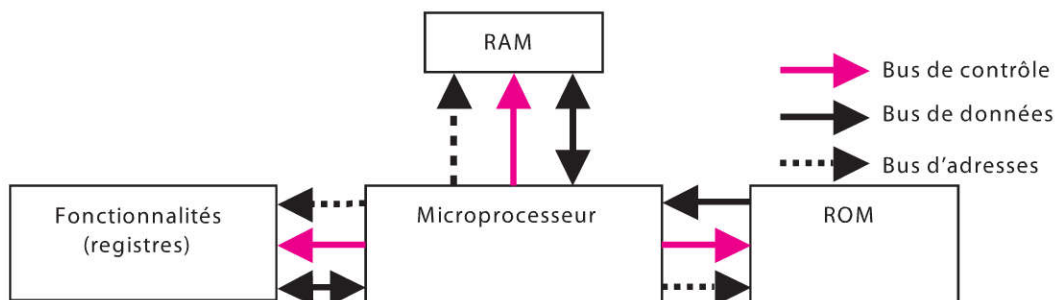
- d'un Microprocesseur BA=A0 à A15 et BD= D0 à D7 ;
- d'une mémoire RAM de 1Koctet (A0 à A9) en haut du plan mémoire ;
- d'une mémoire ROM de 4Koctet (A0 à A11) en bas du plan mémoire ;
- d'un composant d'interface RS232 (A0 à A2 car 8 registres en interne, se situant à partir de l'adresse 0xA0000.

Le plan mémoire équivalent sera représenté par la figure ci-dessous :



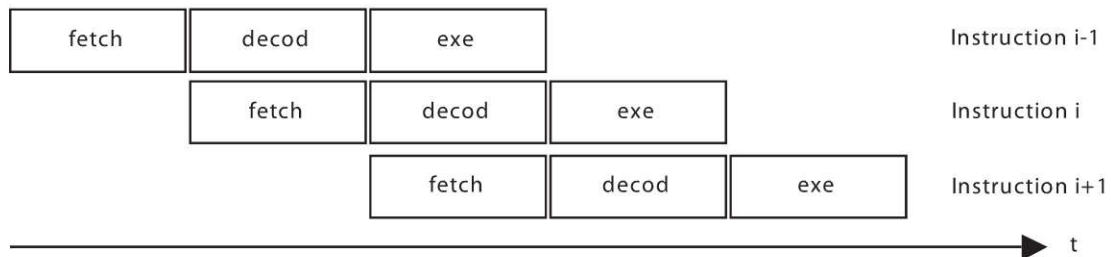
c) L'architecture Harvard

Cette architecture permet de réduire les limitations rencontrées dans l'architecture Von Neumann. Elle repose sur l'utilisation de plusieurs plans mémoires distincts, chaque plan mémoire étant géré par un bus d'adresses, un bus de données et un bus de contrôle spécifique.



Cette approche permet au microprocesseur d'utiliser simultanément ses plans mémoire. Cette architecture est souvent associée à :

- l'utilisation d'un **jeu d'instructions réduit** (RISC – *Reduced Instruction Set Computer*) constitué d'un seul mot mémoire programme (identique à la taille du bus de données de la zone programme). Il ne faut plus qu'un seul accès (1 cycle) pour récupérer l'instruction à exécuter ;
- une organisation en **pipeline** du cycle de lecture (fetch), décodage (decod), exécution (exe), des instructions autorisant une augmentation (théorique) des performances dans un rapport équivalent au nombre de découpage.



3. EN PRATIQUE

Les microprocesseurs, selon le point de vue adopté, font partie du monde de l'informatique ou de celui de l'électronique. L'électronique utilise de plus en plus de microcontrôleurs qui sont un sur-ensemble des processeurs car ils regroupent au sein d'un même composant :

- un microprocesseur ;
- des mémoires de différents types ;
- des fonctionnalités (E/S parallèles, bus de communication, timer, des convertisseurs CAN et CNA, etc.).

Leur mise en œuvre est décrite dans la [fiche 56](#).

61 Chaîne d'acquisition et échantillonnage

Mots clés

Actionneur, capteur, CAN, CNA, échantillonnage, quantification, numérisation, théorème de Shannon, fréquence de Nyquist.

1. EN QUELQUES MOTS

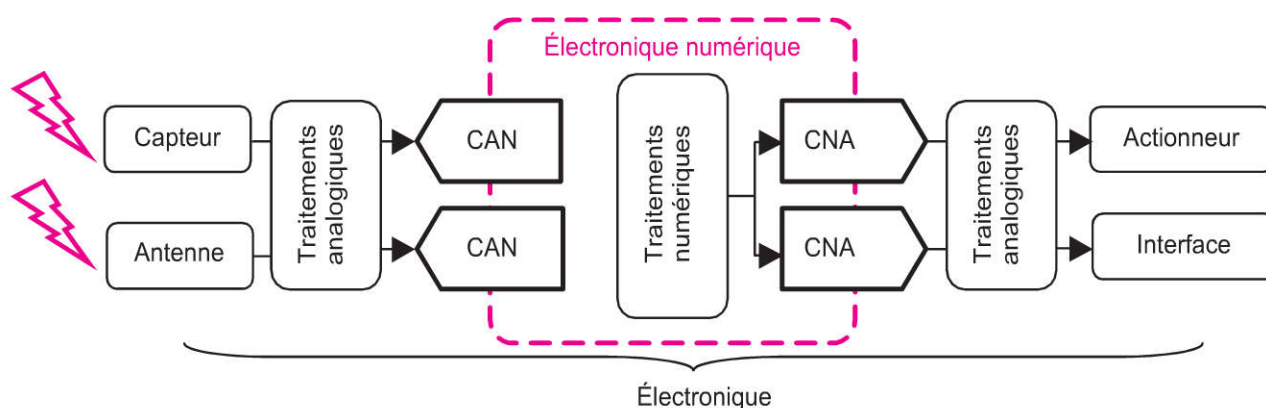
Le monde qui nous entoure est **analogique** : les grandeurs physiques peuvent prendre une **infinité de valeurs continues** dans le temps. D'un autre côté, les **technologies numériques** sont performantes, parfaitement reproductibles et souvent reprogrammables. Cette souplesse est telle que le « numérique » est incontournable dans nos vies quotidiennes. Voyons comment l'électronique permet de concilier ces deux aspects.

2. CE QU'IL FAUT RETENIR

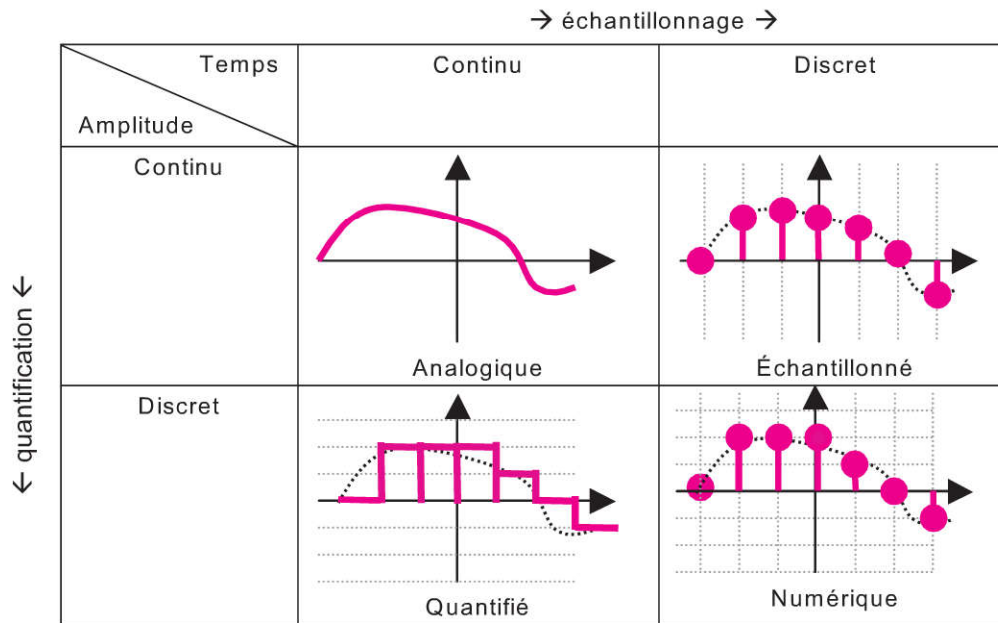
Entre le « monde extérieur » et un système électronique se trouvent d'un côté les **capteurs**, capables de convertir une énergie physique (température, mouvement, pression, électromagnétisme...) en énergie électrique porteuse d'information, et de l'autre, les **actionneurs** et les **interfaces**, conçus pour effectuer l'opération réciproque.

À l'interface entre les domaines analogiques et numériques se trouvent des convertisseurs capables de traduire les données (**convertisseurs analogiques/numériques** – ou CAN – et **convertisseurs numériques/analogiques** – ou CNA).

Une chaîne d'acquisition complète peut se schématiser comme suit :



Le signal physique est capté, transformé en énergie électrique puis traité électroniquement dans le domaine analogique ce qui correspond généralement à de l'amplification (fiches 37, 40, 41 et 47), du filtrage (fiches 26, 27, 50 et 51) et de la démodulation (fiches 66, 67). Une fois mis en forme, le signal électrique analogique est numérisé, ce qui correspond aux opérations d'échantillonnage et de quantification. L'**échantillonnage** consiste à découper périodiquement un signal analogique pour n'en retenir qu'une valeur, soit un **échantillon**, à un instant donné. La **quantification** fait correspondre l'amplitude de cet échantillon à une valeur numérique déterminée d'une échelle donnée. Typiquement, une échelle de quantification comprend 2^N valeurs qui sont, le plus souvent, équiréparties. L'association des opérations d'échantillonnage et de quantification s'appelle la numérisation. En d'autres termes, un signal numérique est issu de l'échantillonnage et de la quantification d'un signal analogique.



Le signal numérique peut être manipulé de façon parfaitement reproductible. Les traitements numériques du signal les plus courants sont le filtrage, la mémorisation, la FFT (*fast Fourier Transform*), le codage et le décodage, la démodulation et la modulation... Les algorithmes de traitements numériques du signal sont très nombreux et doivent faire l'objet d'une étude spécifique allant au-delà du cadre des bases de l'électronique.

Toutes les opérations décrites ci-dessus ont leur réciproque et il est ainsi possible de commander des interfaces et des actionneurs à partir de données numériques.

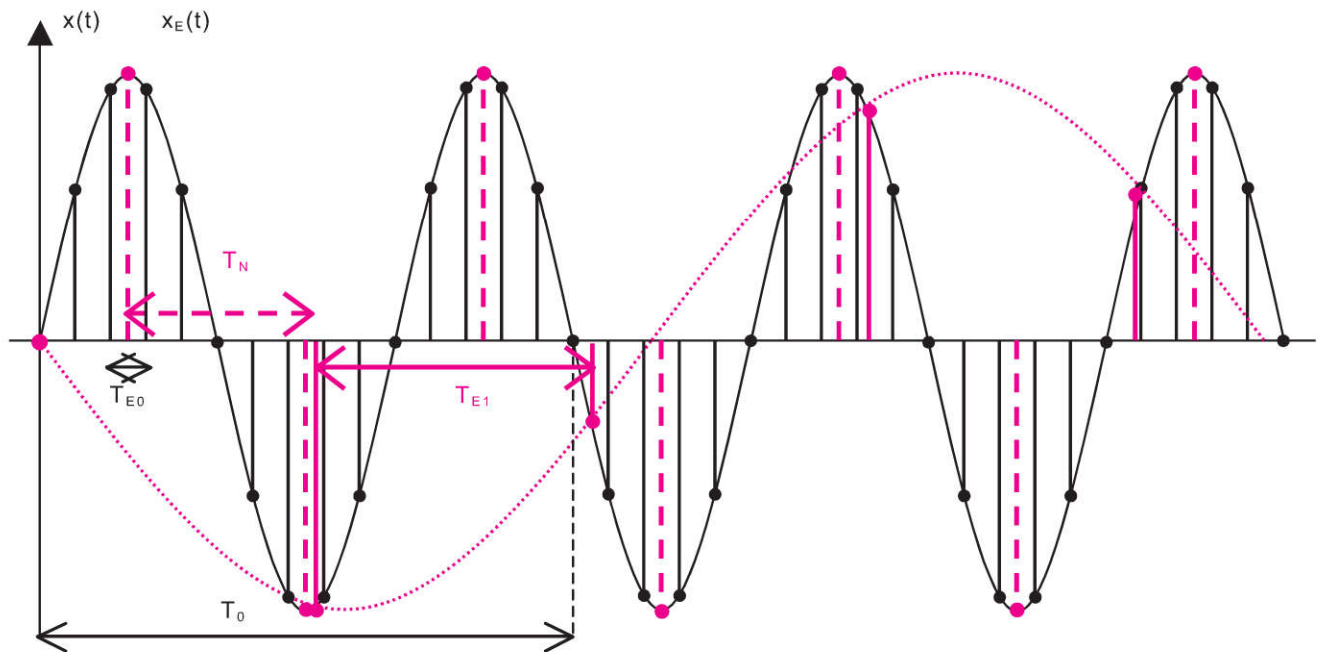
3. EN PRATIQUE

a) Théorème de Shannon

L'échantillonnage consiste à capturer périodiquement des valeurs d'un signal analogique. Si la période choisie est faible par rapport à la vitesse de variation du signal, certains événements, comme des changements brutaux, des phénomènes apparaissant à des fréquences élevées, ne seront pas ou seront mal acquis. Il faut donc que la fréquence d'échantillonnage soit suffisante par rapport au signal analogique, mais pas démesurément grande afin de ne pas gaspiller des ressources électroniques matérielles (mémoires...) et du temps de calculs logiciels.

La limite, appelée fréquence de Nyquist (F_N), est donnée par le théorème de Shannon : la fréquence d'échantillonnage F_e doit être au minimum le double de la fréquence maximale $F_{\max}$ comprise dans le signal échantillonné.

La figure suivante représente en trait plein noir le signal analogique sinusoïdal $s(t)$ de période T_0 à échantillonner. Un échantillonnage consistant à prélever 10 échantillons par période (période d'échantillonnage T_{E0}) reproduit correctement le signal d'origine. Lorsque celui-ci est inférieur à deux échantillons par période (échantillons en traits pleins rouges, période d'échantillonnage T_{E1}), nous constatons qu'il peut y avoir une ambiguïté sur la représentation du signal d'origine : c'est la sinusoïde représentée en pointillés rouges qui peut être reconstituée à partir de ces échantillons. Le cas limite est celui pour lequel deux échantillons (en pointillés rouges) sont prélevés par période T_0 : $T_N = T_0 / 2$.



Le théorème de Shannon se traduit donc par les relations suivantes :

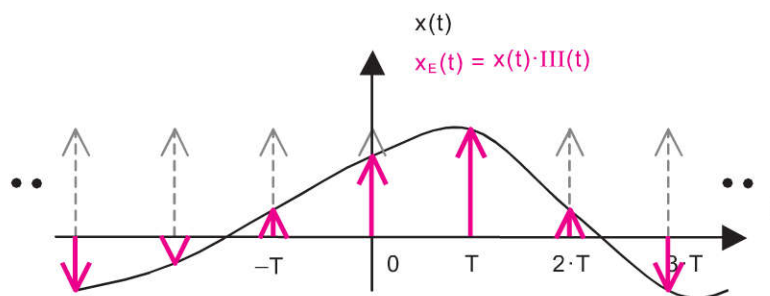
$$f_N = 2 \cdot f_{\max} \Rightarrow f_E \geq 2 \cdot f_{\max}$$

Avec $f_N = \frac{1}{T_N}$, $f_E = \frac{1}{T_E}$ et $f_{\max} = \frac{1}{T_{\max}}$, ce théorème peut être réécrit en raisonnant en termes de périodes :

$$T_N = \frac{T_{\max}}{2} \Rightarrow T_E \leq \frac{T_{\max}}{2}$$

Sous un point de vue mathématique, l'échantillonnage consiste à multiplier le signal continu par un peigne de Dirac (fiche 14).

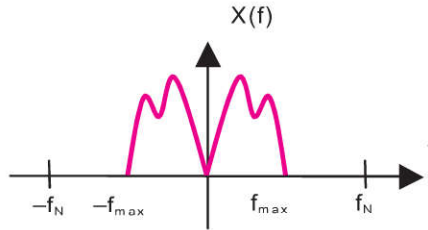
$$\delta(t) \cdot \text{III}_T(t) = x(t) \cdot \sum_{k=-\infty}^{+\infty} \delta(t - k \cdot T) = \sum_{k=-\infty}^{+\infty} x(k \cdot T) \cdot \delta(t - k \cdot T)$$



Représentation spectrale de l'échantillonnage

b) Représentation spectrale de l'échantillonnage

Soit $x(t)$ un signal quelconque de fréquence maximale $f_{\max}$ et $X(f)$ sa transformée de Fourier (fiche 16).

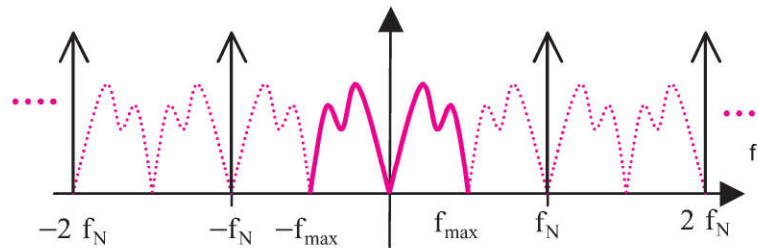


Dans le domaine temporel, échantillonner un signal revient à multiplier celui-ci par un peigne de Dirac de période T_E .

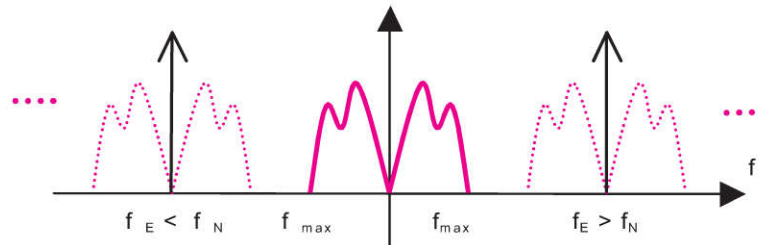
La transformée de Fourier d'un signal $x(t)$ échantillonné à la période T correspond au produit de convolution de la transformée de Fourier $X(f)$ de $x(t)$ et d'un peigne de Dirac.

$$x(t) \cdot \text{III}_{T_E}(t) \xrightarrow{\text{TF}} \frac{1}{T_E} \cdot X(f) * \frac{1}{T_E}(f)$$

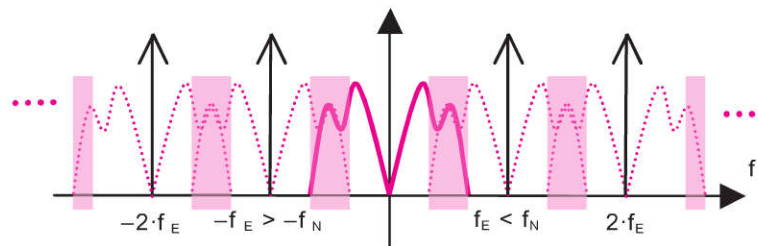
Dans le cas limite où $f_E = f_N = 2 f_{\max}$, les spectres répliqués du signal $X(f)$ se retrouvent positionnés bord à bord.



Si la fréquence d'échantillonnage est plus élevée que f_N , les éléments issus de la réplication de $X(f)$ s'éloignent.



Réciproquement, si la fréquence d'échantillonnage ne respecte pas le théorème de Shannon, les éléments issus de la réplication de $X(f)$ se recouvrent : le spectre résultant ne représente plus le signal d'origine, c'est le phénomène illustré dans le domaine temporel sur le graphique introduisant le théorème de Shannon.



Attention

Il ne faut pas confondre produit de convolution par un Dirac et multiplication par un Dirac (fiche 18).

62 Convertisseurs Numérique-Analogique

Mots clés

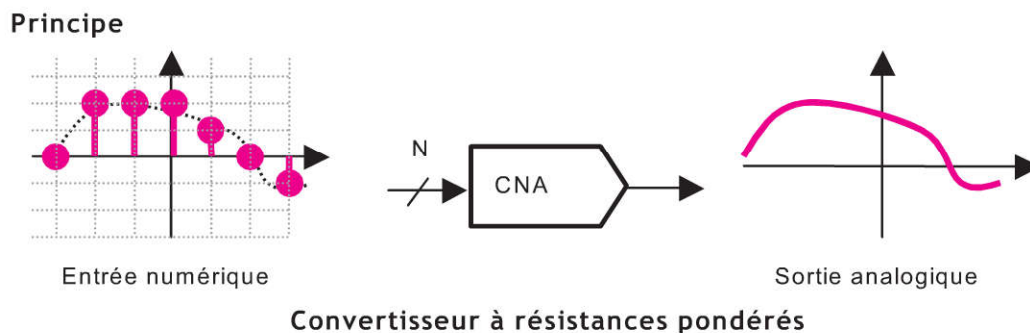
Filtre passe-haut, filtre passe-bas, diagramme de Bode, transmittance, atténuation.

1. EN QUELQUES MOTS

Le rôle d'un convertisseur numérique-analogique (CNA, ou DAC pour *Digital-to-Analog Converter*) est de générer un signal analogique à partir d'un signal numérique codé sur n bits (fiches 53 et 61). Les deux architectures les plus répandues sont le CNA à résistances pondérées et le CNA à réseau R-2R.

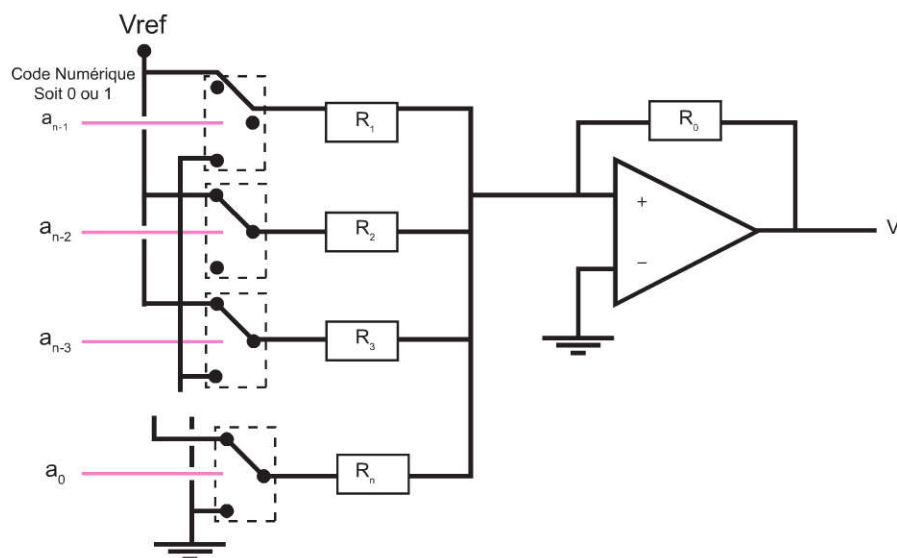
2. CE QU'IL FAUT RETENIR

a) Principe



b) Convertisseur à résistances pondérées

Ce CNA est basé sur le principe d'un additionneur à amplificateur opérationnel (fiche 49) : la tension d'entrée est pondérée par des résistances dont les valeurs sont étagées par puissances de 2.



Si le bit i est à 1, le commutateur correspondant se met en position de connecter la résistance $n-i$ à la tension de référence V_{ref} . Dans le cas contraire, le commutateur se met à la masse (0 V).

$$R_1 = 2^0 \cdot R_1, R_2 = 2^1 \cdot R_1, R_3 = 2^2 \cdot R_1, R_n = 2^{n-1} \cdot R_1$$

$$V_o = -V_{ref} \cdot \frac{R_o}{R_1} \cdot \left(\frac{a_{n-1}}{2^0} + \frac{a_{n-2}}{2^1} + \frac{a_{n-3}}{2^2} + \dots + \frac{a_0}{2^{n-1}} \right)$$

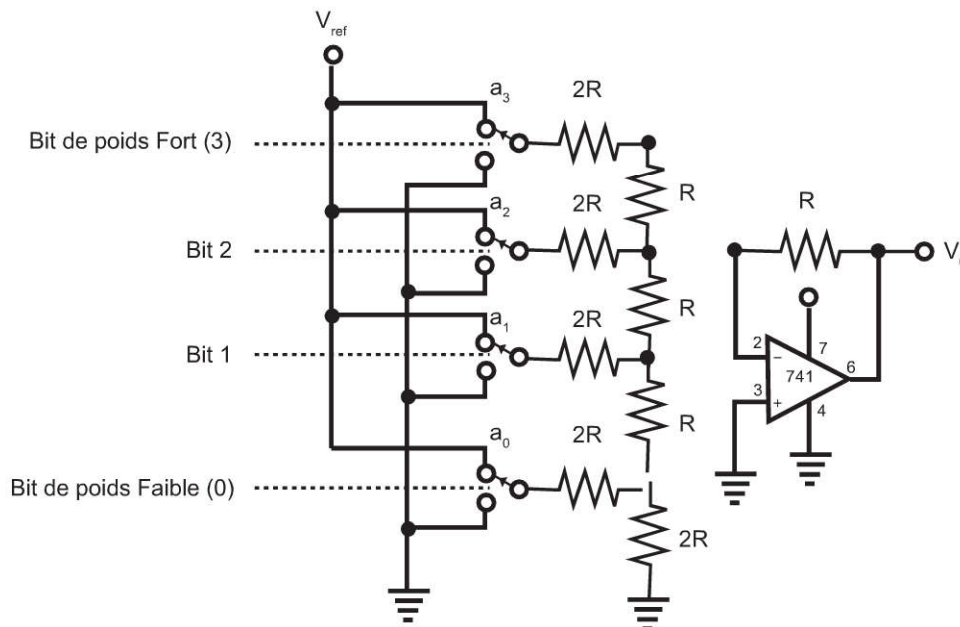
Dans un CNA, l'écart entre deux valeurs analogiques contiguës (quantum ou pas de quantification) est : $\varepsilon = \frac{R_0 \cdot V_{ref}}{2^{n-1} \cdot R_1}$

La dynamique de sortie, c'est-à-dire la différence entre la plus petite et la plus grande valeur analogique générée est : $V_{\max} = \frac{2 \cdot R_0 \cdot V_{ref}}{R_1}$

c) Convertisseur à réseau R-2R

Ce CNA est basé sur le principe d'un montage inverseur à amplificateur opérationnel (fiche 50) : la tension d'entrée est pondérée par des résistances de valeurs différentes.

L'exemple (schéma) montre un CNA à réseau R-2R sur 4-bits.



si pour $n^{\text{ème}}$ bits est «1» si elle est reliée à une tension V_{ref} et «0» si elle est reliée à la masse, la tension de sortie V_0 peut donc être calculée comme suit :

$$V_0 = \frac{V_{ref}}{2^N} \cdot (a_{N-1} \cdot 2^{N-1} + a_{N-2} \cdot 2^{N-2} + a_{N-3} \cdot 2^{N-3} + \dots + a_1 \cdot 2^1 + a_0 \cdot 2^0)$$

3. EN PRATIQUE

Le CNA à résistances pondérées utilise des résistances de valeurs très différentes, avec un rapport de 2^{n-1} entre la plus grande et la plus faible. En sachant que $R > 5 \text{ k}\Omega$, cela pose des problèmes de précision des éléments résistifs et des difficultés d'intégration. Cette architecture est généralement écartée au profit du réseau R-2R ne nécessitant que deux valeurs de résistances. Ce deuxième montage permet donc un choix plus aisé des résistances et donc une meilleure précision.

63 Conversion Analogique-Numérique

Mots clés

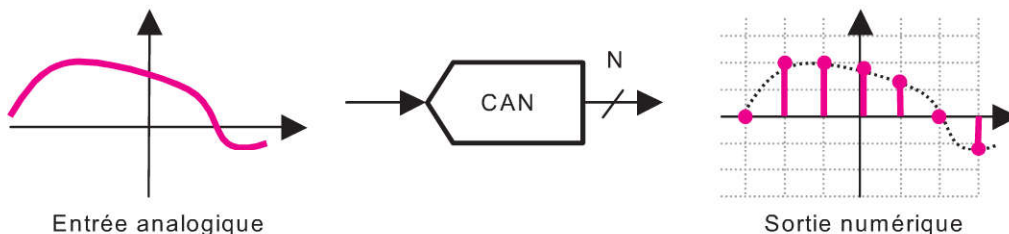
CAN à approximations successives, flash, pipeline, sigma-delta, résolution, pas de quantification, SNR, ENOB, SFDR.

1. EN QUELQUES MOTS

Un convertisseur analogique-numérique (noté CAN) est un dispositif électronique assurant les fonctions d'échantillonnage et de quantification d'un signal analogique (fiche 61). Ces composants fonctionnent selon des principes différents et offrent des performances variées. Les paramètres permettant d'évaluer ces performances sont décrits ici.

2. CE QU'IL FAUT RETENIR

a) Principe



b) Nombre de bits, résolution, pas de quantification, LSB

Le nombre de bits d'un CAN correspond à la taille des données numériques N codant le signal analogique d'entrée. Les valeurs typiques de N sont comprises entre 6 et 24 bits. Le nombre de paliers de conversion correspond au nombre de codes numériques différents que peut produire le CAN à sa sortie, il est donc de 2^N .

La résolution, ou pas de quantification, indique la plus petite valeur de tension à laquelle correspond un code numérique. Si la plage de tension permise en entrée est notée V, la résolution ϵ du

$$\text{CAN est donc } \epsilon = \frac{V}{2^N}$$

Cette valeur correspond à la marge d'erreur avec laquelle le CAN code une tension d'entrée comprise dans la plage de tension V. Elle correspond au poids d'un bit de poids faible (LSB : *Least Significant Bit*).

c) Fréquence d'échantillonnage

Les fréquences d'échantillonnage sont exprimées en nombre d'échantillons par seconde. Les valeurs sont généralement exprimées en ksp/s,Msp/s ou Gsp/s (kilo/Méga/Giga Samples Per Second – kilo/Méga/Giga échantillons par seconde).

d) Bruit de quantification et SNR

On remarque que le pas de quantification ϵ présente une densité de probabilité uniforme dans l'intervalle $[-q/2, q/2]$. En normalisant la probabilité de ϵ , on obtient :

$$p(\epsilon) = 1/q \text{ pour } -q/2 < \epsilon < q/2$$

$$p(\epsilon) = 0 \text{ ailleurs}$$

La moyenne quadratique de ε ou puissance de bruit s'écrit :

$$E(\varepsilon^2) = \frac{1}{q} \int_{-q/2}^{q/2} \varepsilon^2 d\varepsilon = \frac{q^2}{12}$$

La puissance du signal sinusoïdal d'amplitude A est égale à $A^2/2$ (fiche 8 et 12). Pour un signal pleine échelle (c'est-à-dire pour $A = V$), on obtient :

$$q = \frac{A}{2^{N-1}}$$

Déduisons-en le rapport signal à bruit linéaire et en dB (fiche 19) :

$$\text{SNR}_{\text{lin}} = \frac{A^2}{q^2/12} = 3 \cdot 2^{2 \cdot N-1}$$

$$\text{SNR}_{\text{dB}} = 10 \cdot \log(3 \cdot 2^{2 \cdot N-1}) \approx 6,02 \cdot N + 1,76$$

où N représente le nombre de bits du CAN considéré.

Les pas de quantification n'étant pas parfaitement identiques dans un CAN réel, le SNDR (*Signal-to-Noise and Distorsion Ratio*) inclut cette non-uniformité

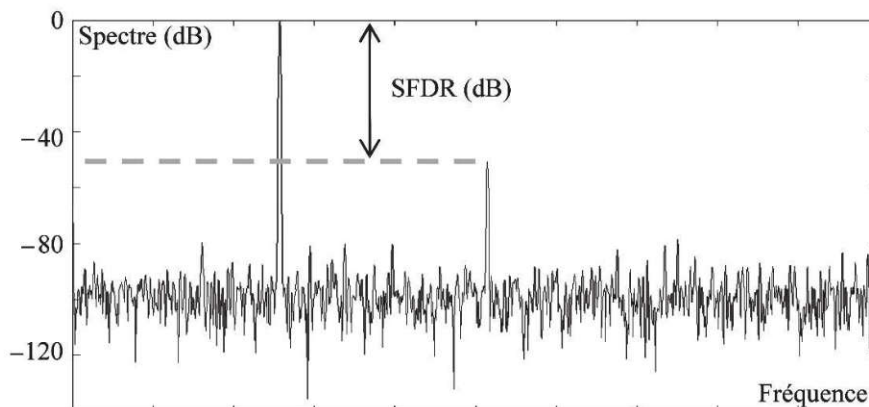
e) ENOB

L'ENOB (*Effective Number Of Bits* – nombre de bits effectifs) reflète l'écart entre le niveau de bruit mesuré en sortie d'un CAN donné et le niveau de bruit théorique d'un CAN idéal numérisant le signal sur le même nombre de bits N .

$$\text{ENOB} = \frac{\text{SNDR} - 1,76}{6,02}$$

f) SFDR

Le SFDR (*Spurious Free Dynamic Range* – dynamique sans fréquence parasite) correspond à la différence entre la puissance d'une sinusoïde numérisée en pleine échelle et la puissance de la fréquence parasite la plus élevée.



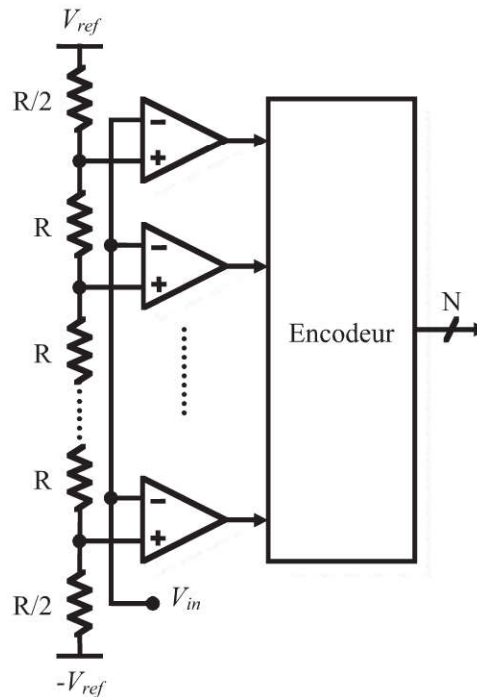
3. EN PRATIQUE

De nombreuses familles de CAN sont disponibles sur le marché. Les plus courantes sont présentées ici.

a) Flash

Le convertisseur flash est une architecture de CAN dont le principe de fonctionnement est très intuitif. Ce circuit utilise $2^N - 1$ comparateurs pour fournir un résultat numérique sur N bits. Les

comparateurs sont basés sur l'utilisation d'amplificateurs opérationnels (fiches 48 et 49) et d'un réseau de résistances faisant office de diviseur de tension (fiche 23).

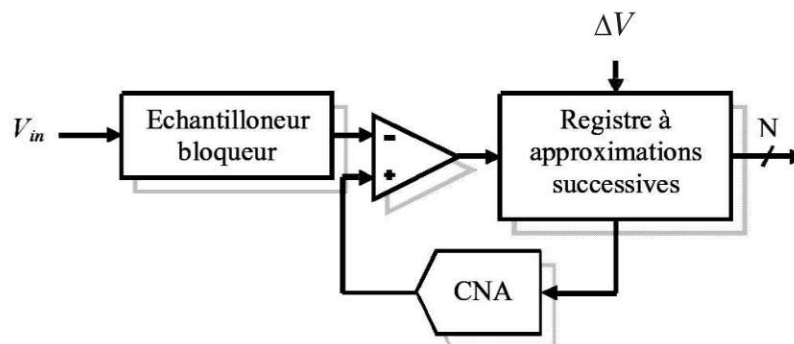


Les seuils des comparateurs sont fixés par le réseau de résistances pour qu'ils soient équirépartis et distants d'un LSB. La quantification s'effectue en une première étape, la seconde consistant à traduire le code thermométrique issu du réseau de comparateurs en un code binaire à l'aide d'un circuit d'encodage numérique.

L'avantage majeur de ce CAN est sa capacité à travailler à des fréquences très élevées, éventuellement supérieures à 1 Gsps. L'inconvénient principal provient du réseau de comparateurs dont la taille est doublée pour chaque bit supplémentaire. Les résolutions proposées avec les CAN flash sont donc en général de l'ordre de 8 bits afin de limiter la consommation, la taille du circuit et les erreurs de seuil dues à la difficulté de produire un réseau de comparateurs de grande taille et uniforme.

b) CAN à approximations successives

Contrairement au CAN flash où l'opération de conversion est totalement parallélisée, ce CAN est un système bouclé contenant un CNA dans sa boucle de retour.

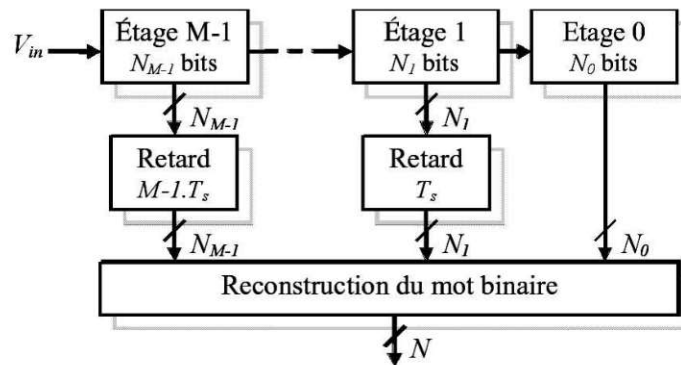


Le rôle du registre à approximations successives (SAR) est de présenter en entrée du CNA successivement les différents bits en commençant par celui de poids le plus fort. La plage de recherche de la valeur à numériser étant divisée par 2 pour chaque nouveau bit à évaluer, ce CAN utilise le principe de la dichotomie.

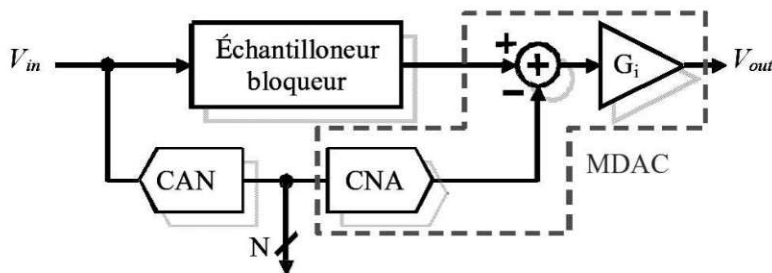
Pour une résolution de N bits, la conversion s'effectue en N périodes d'horloge. Ces CAN offrent des capacités d'intégration élevées puisqu'ils ne contiennent qu'un comparateur. Ce principe ne permet pas d'obtenir des fréquences de fonctionnement élevées (inférieures à 1 MSPS), mais il représente un exemple de sérialisation de l'opération de conversion.

c) Pipeline

Les CAN pipeline ont la particularité de répartir la quantification le long d'une chaîne de convertisseurs. Au niveau de la parallélisation des opérations, le CAN pipeline est un compromis entre les deux architectures précédemment indiquées. En effet, alors que le CAN flash requiert 2^N comparateurs et que le CAN à approximations successives n'en contient qu'un seul, le CAN pipeline représente un intermédiaire en répartissant l'effort de conversion sur M étages.

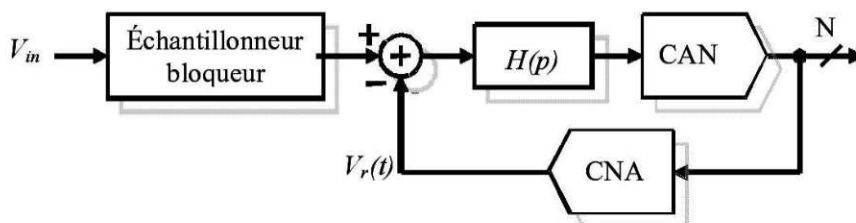


Chacun des M étages est un bloc de conversion élémentaire qui numérise le signal présenté à son entrée et fournit à l'étage suivant le signal d'erreur analogique (reste des opérations antérieures).



d) CAN Sigma-Delta

La structure Sigma-Delta, ou $\Sigma\Delta$, fait partie de la famille des CAN à suréchantillonnage. Elle est constituée d'une boucle appelée modulateur sigma delta et d'un filtre numérique de réponse $H(p)$.



Le cas le plus commun consiste à utiliser un CAN 1 bit (simple comparateur), ce qui rend inutile le CNA en rétroaction, et un simple intégrateur pour $H(z)$. La sortie est donc constituée d'une trame de bits traitée par un filtre décimateur pour obtenir une image numérique de l'entrée à la fréquence de Nyquist.

Cette architecture, utilisée en audio, est très performante jusqu'à une fréquence d'utilisation de plusieurs Msps en permettant d'atteindre des résolutions élevées allant jusqu'à 24 bits.

64 Boucle à verrouillage de phase

Mots clés

Asservissement, PLL, VCO, filtrage, diviseur de fréquence, comparateur de phase.

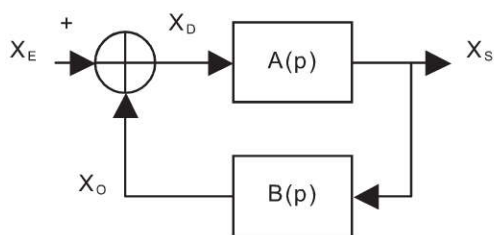
1. EN QUELQUES MOTS

Les boucles à verrouillage de phase ou PLL (*Phase-Locked Loop*) sont des systèmes synthétiseurs de fréquence qui permettent d'obtenir des fréquences stables. C'est un système asservi largement utilisé en transmission numérique et dans les systèmes de démodulation de signaux modulés en fréquences (fiches 66 et 67).

2. CE QU'IL FAUT RETENIR

a) Systèmes asservis : rappels

Dans le cas général, un système asservi se présente sous la forme suivante :

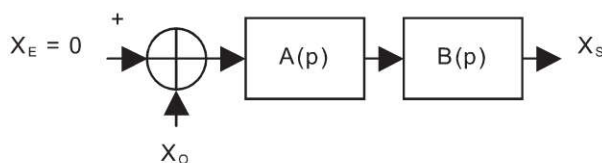


Élément de système	Équation
Réseau direct	$X_S = A(p) \cdot X_D$
Réseau de retour	$X_O = B(p) \cdot X_S$
Soustracteur	$X_D = X_E - X_O$

Sa fonction de transfert, dans le domaine de Laplace, s'écrit sous la forme :

$$H(p) = \frac{X_S}{X_E} = \frac{A(p)}{1 + A(p) \cdot B(p)}$$

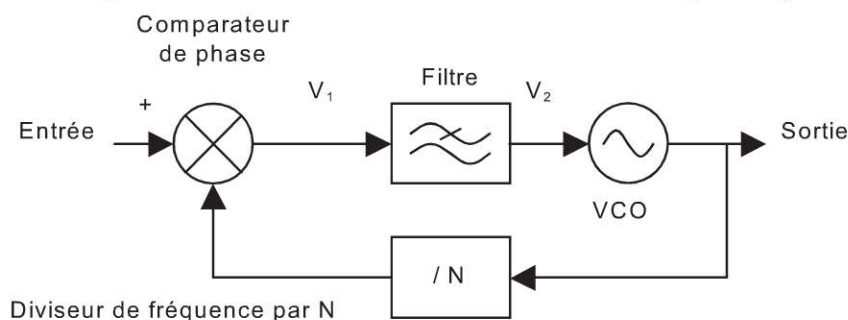
Le gain en boucle ouverte $G(p)$ s'obtient en ouvrant le système au niveau de la grandeur X_O , avec X_O l'entrée du système et X_S sa sortie.



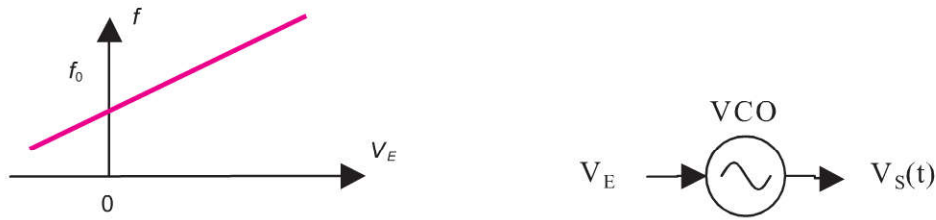
$$G(p) = \frac{X_S}{X_O} = -A(p) \cdot B(p)$$

b) PLL : principe

Une boucle à verrouillage de phase est constituée d'un oscillateur contrôlé en tension, d'un filtre, d'un comparateur de phase et éventuellement d'un diviseur de fréquence par N.



L'oscillateur contrôlé en tension (ou VCO pour *Voltage-Controlled Oscillator*) est un dispositif générant un signal sinusoïdal dont la fréquence est proportionnelle à la tension d'entrée.



Dans le cas idéal, la fréquence de sortie se calcule grâce à une fonction affine. Nous pouvons en déduire l'expression de la phase instantanée (après transformée de Laplace) sachant que la fréquence est la dérivée de la phase.

$$f = f_0 + K_{VCO} \cdot V_E \Rightarrow \varphi(p) = \frac{K_{VCO} \cdot V_E(p)}{p}$$

K_{VCO} s'exprime en $\text{Hz} \cdot \text{V}^{-1}$

Le signal de sortie du VCO est une sinusoïde dont la fréquence dépend de K_{VCO} :

$$V_S(t) = A \cdot \sin[2 \cdot \pi \cdot (f_0 + K_{VCO} \cdot V_E) \cdot t + \varphi_0]$$

Le comparateur de phase génère une tension continue proportionnelle à la différence de phase entre les deux signaux d'entrée :

$$V_\varphi(t) = K_\varphi \cdot [\varphi_1(t) - \varphi_2(t)]$$

Le filtre de boucle, de fonction de transfert $F(p)$, a pour rôle de stabiliser le système.

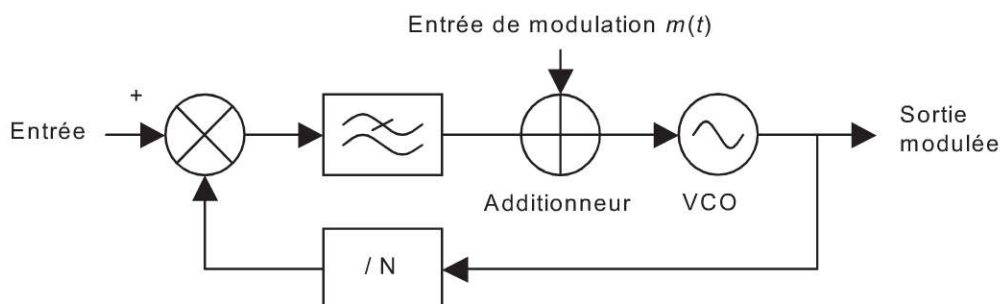
Le diviseur de fréquence par N est tel que :

$$V_S(t) = A \cdot \sin\left(\frac{\omega}{N} \cdot t\right)$$

La fonction de transfert s'écrit alors : $H_{PLL}(p) = \frac{\varphi_S(p)}{N \cdot \varphi_E(p)} = \frac{K_{VCO} \cdot K_\varphi \cdot F(p)}{N \cdot p + K_{VCO} \cdot K_\varphi \cdot F(p)}$

3. EN PRATIQUE

Les PLL sont très utilisées pour les modulations et démodulations de fréquence (fiches 66 et 67). La démodulation consiste à mettre un signal FM sur l'entrée, et le signal démodulé se trouve en sortie du VCO. Pour la modulation, il suffit de rajouter le signal modulant dans la boucle avec un additionneur.



65 Transmission de l'information

Mots clés

Modulation démodulation, émetteur, récepteur, source, canal.

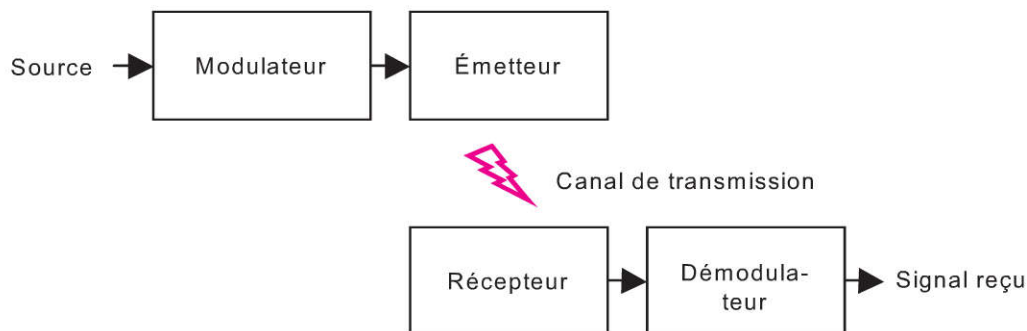
1. EN QUELQUES MOTS

Les télécommunications consistent à **transmettre** et à **recevoir des informations à distance** (téléphonie, Internet, Wifi, Bluetooth, liaisons satellites, TNT...). La transmission d'informations suppose :

- d'avoir des informations à transmettre ;
- d'adapter la forme de ces informations au milieu dans lequel ces informations vont voyager (modulation) ;
- que le récepteur connaisse la technique de modulation afin de pouvoir extraire l'information du signal reçu.

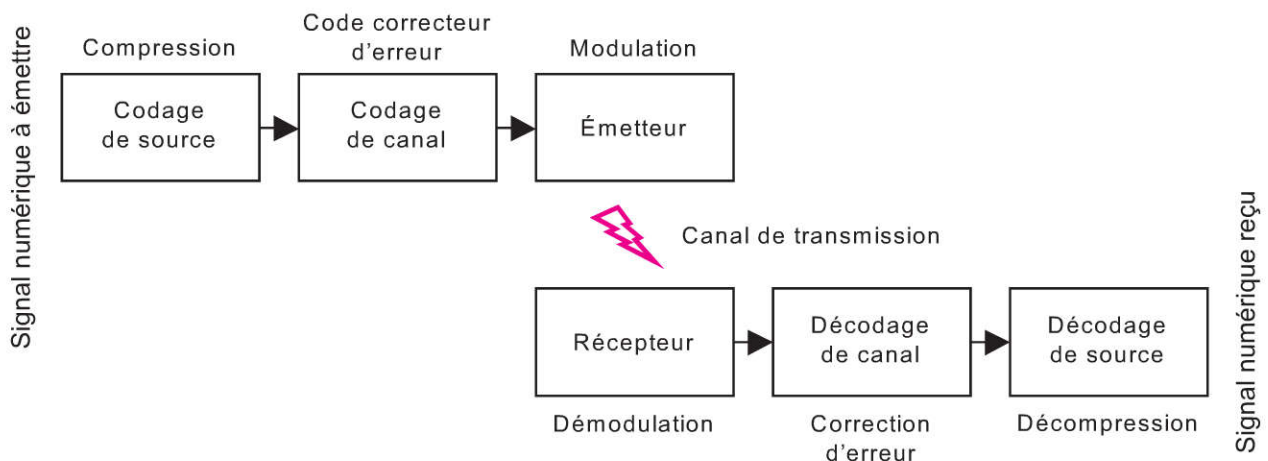
Les informations à transmettre sont sous **forme analogique ou numérique**. Selon la forme, les techniques de transmission de l'information diffèrent.

2. CE QU'IL FAUT RETENIR



Si la transmission est parfaite, le signal reçu est identique à la source. Dans le cas d'un signal analogique, la source est forcément perturbée par le **canal de transmission** et également par le bruit induit naturellement par les systèmes électroniques mis en œuvre.

Différentes **modulations analogiques** de base sont présentée dans la [fiche 66](#).



Dans le cas d'une **information numérique**, les sources de perturbation sont les mêmes, mais du fait du caractère binaire de l'information, la distinction entre un 1 et un 0 étant *a priori* plus nette, on peut espérer extraire du signal reçu une copie parfaite de l'information émise.

Le **codage de source** supprime les éléments binaires qui ne sont pas utiles à l'information. Le **décodage de source** effectue l'opération réciproque. Les fonctions de **compression** et de **décompression** sont des exemples de codage et de décodage de source (compression audio ou vidéo par exemple).

Le **codage de canal** rajoute des informations au signal afin de préserver l'immunité de l'information : il permet d'évaluer la validité de l'information reçue et éventuellement de détecter les erreurs dues à un bruit trop élevé, à une déformation du signal trop importante (multitraitement dû à des réflexions sur des bâtiments par exemple) ou à la faiblesse du signal reçu (atténuation due à la distance, aux conditions atmosphériques...).

L'**émetteur** se charge de conditionner le signal pour qu'il soit compatible avec son déplacement dans le canal de transmission. Le signal est **modulé, filtré** (pour éviter d'encombrer le spectre et de perturber d'autres transmissions), **amplifié** pour palier à l'atténuation et **émis** par antenne, laser, port de communication... Le signal est ensuite **capté par le récepteur** (antenne, photodiode, port de communication...) qui effectue les opérations de **filtrage** (pour ne retenir que l'information transportée dans le canal souhaité), d'**amplification** pour que le signal ait une amplitude suffisante pour être traité par le système électronique en amont et de **démodulation** pour extraire l'information pertinente.

3. EN PRATIQUE

La performance d'une transmission est évaluée afin d'estimer la validité de l'information reçue par rapport à l'information envoyée. Bien sûr, le récepteur n'a accès à cette dernière que *via* le canal de communication et ne peut pas faire de comparaison directe avec la source.

Dans le **domaine analogique**, la transmission est évaluée en calculant le rapport signal à bruit (SNR, [fiche 68](#)).

Pour les **transmissions numériques**, c'est le taux d'erreur binaire (TEB ou BER pour *Binary Error Rate*) qui sert à évaluer la qualité de la transmission :

$$\text{TEB} = \frac{N_{\text{bit erronés}}}{N_{\text{bit transmis}}}$$

L'efficacité spectrale indique le rapport entre le débit binaire atteint et l'espace fréquentiel occupé par le signal modulé. Cette grandeur est associée à la notion de rendement : il est souhaitable de transmettre avec un débit élevé dans un canal de faible encombrement spectral.

$$\eta = \frac{\text{débit}}{\text{bande occupée}} = \frac{D}{B}$$

η peut être noté en bits/s/Hz.

66 Modulations analogiques

Mots clés

Modulation d'amplitude, de fréquence, de phase, signal modulant, modulé, bande de base, indice de modulation.

1. EN QUELQUES MOTS

Le **signal en bande de base** désigne le signal avant modulation ou après démodulation. Il **contient l'information** : il occupe une bande de fréquence comprise entre 0 et $f_{\max}$. En général, on considère que la fréquence maximale correspond à une puissance moitié (-3 dB) de la puissance maximale. L'objectif de la modulation est de **transposer le signal** en bande de base vers des fréquences plus élevées afin de l'adapter au canal de transmission (câble coaxial, ondes radio...). C'est le cas par exemple pour la radio commerciale FM qui occupe la bande de fréquences 88-108 MHz, mais qui transporte des informations audio (0 à 20 kHz environ) en stéréo.

Les techniques permettant de combiner un signal **modulant** (information) et une fréquence porteuse sont classées ici en trois catégories : les **modulations d'amplitude** et les modulations angulaires comprenant la **modulation de phase** et la **modulation de fréquence**.

2. CE QU'IL FAUT RETENIR

a) Modulation d'amplitude

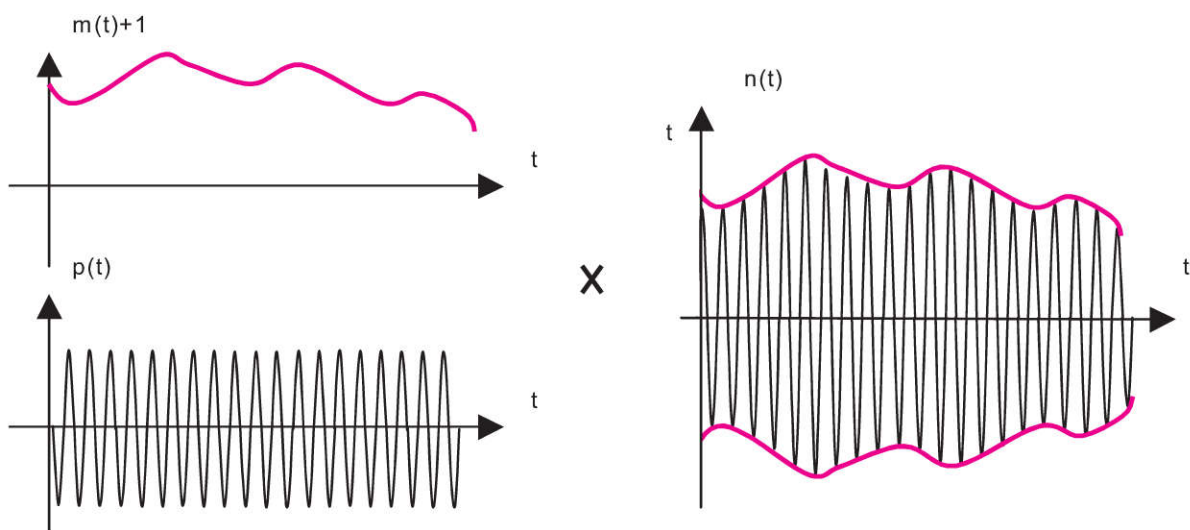
Soient le signal modulant $m(t)$ et la **porteuse** de fréquence beaucoup plus élevée que $f_{\max}$ $p(t) = A \cdot \cos(\omega \cdot t)$. Les modulations d'amplitude produisent un signal $n(t)$ tel que :

$$n(t) = A(t) \cdot \cos(\omega \cdot t)$$

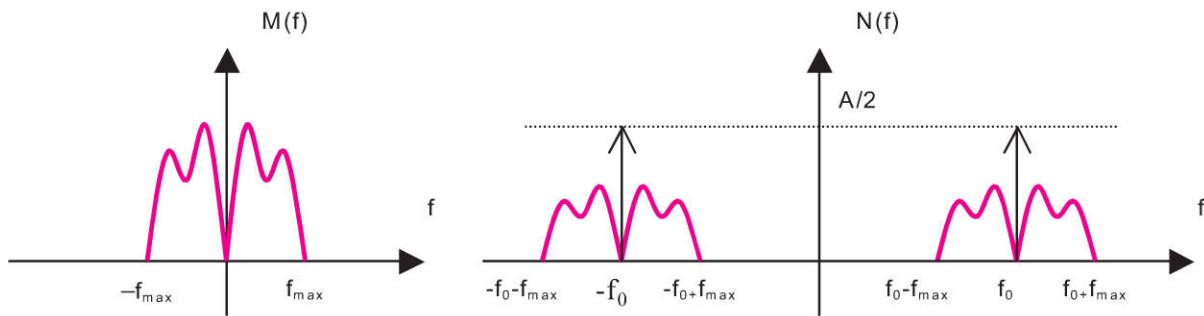
L'amplitude du signal produit dépend du temps. $A(t)$ est fonction du signal modulant $m(t)$.

La modulation d'amplitude double bande consiste à modifier l'amplitude de $p(t)$ avec le signal porteur d'information $m(t)$. Cette opération se fait par multiplication des signaux $m(t)$ et $p(t)$.

$$n(t) = (A + m(t)) \cdot \cos(\omega \cdot t)$$



Soient $M(f)$ et $N(f)$ les résultats des transformées de Fourier des signaux temporels $m(t)$ et $n(t)$. La porteuse sinusoïdale pure $p(t)$ est un Dirac après transformée de Fourier. La multiplication de la modulation d'amplitude dans le domaine temporel devient une convolution dans le domaine fréquentiel (fiche 18).



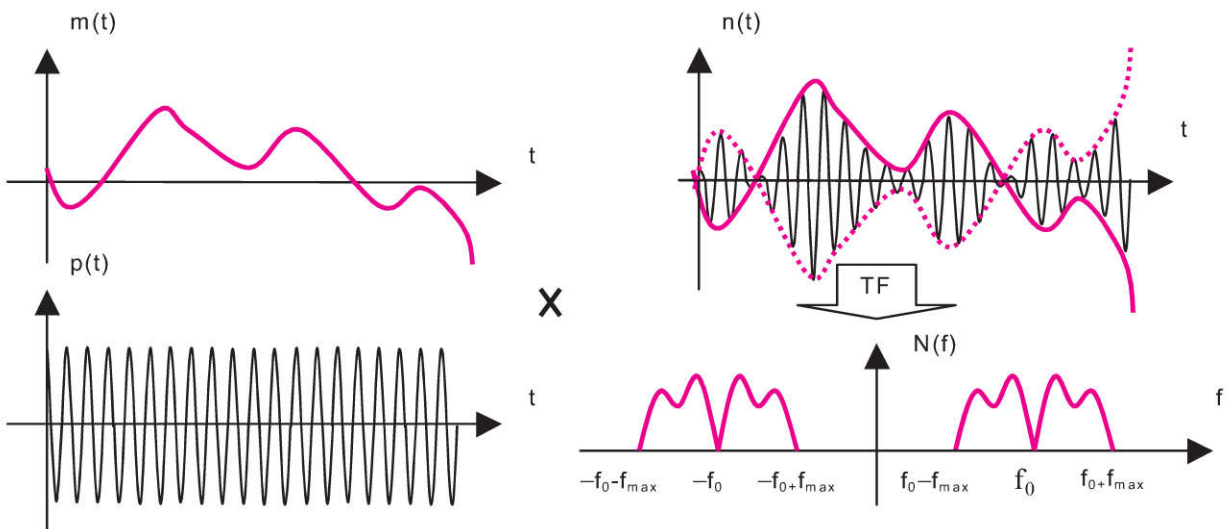
Dans le cas où le signal bande de base est une sinusoïde, prenons $m(t) = B \cdot \cos(\omega_1 \cdot t)$, le signal modulé devient :

$$n(t) = A \cdot (1 + m_A \cdot \cos(\omega_1 \cdot t)) \cdot \cos(\omega \cdot t)$$

avec $m_A = B/A$ l'indice de modulation d'amplitude.

Dans ce cas, le spectre de $M(f)$ est une paire de Dirac de poids $B/2$, et le spectre du signal modulé est le résultat de la convolution de ces Dirac par les Dirac de la porteuse.

La **modulation d'amplitude à porteuse supprimée** est telle que le signal modulant est centré : sa composante continue est supprimée.

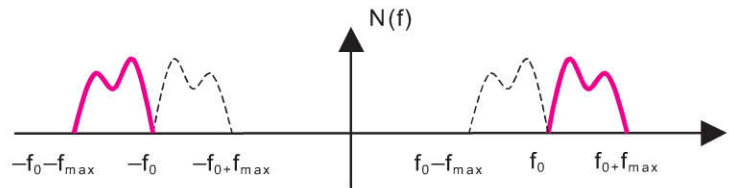


Dans le cas d'un signal modulant $m(t)$ sinusoïdal, le signal modulé s'écrit :

$$n(t) = p(t) \cdot m(t) = A \cdot \cos(\omega \cdot t) \cdot B \cdot \cos(\omega_1 \cdot t) = \frac{A \cdot B}{2} \cdot [\cos((\omega + \omega_1) \cdot t) + \cos((\omega - \omega_1) \cdot t)]$$

Nous voyons apparaître deux composantes sinusoïdales : l'une à $\omega + \omega_1$ et l'autre à $\omega - \omega_1$, mais la porteuse à la pulsation ω n'est pas présente. Cette technique permet de réduire la puissance émise : seul le signal porteur d'information est émis.

La **modulation d'amplitude à bande latérale unique** (BLU) permet encore d'augmenter l'efficacité de la transmission en supprimant la redondance d'information due à la symétrie du spectre du signal en bande de base.



Dans la représentation fréquentielle proposée, c'est la bande latérale inférieure qui est supprimée : elle est retirée du spectre grâce à un filtre à « flancs raide », c'est-à-dire dont l'atténuation à partir de la fréquence de coupure suit une pente quasi verticale.

b) Modulations angulaires

Les modulations angulaires agissent sur la phase instantanée du signal porteur :

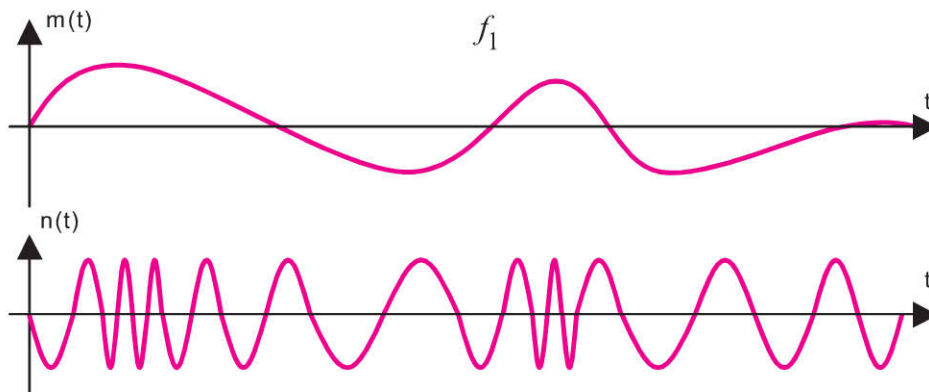
- les **modulations de phase** modifient la valeur de φ ;
- les **modulations de fréquence** modifient la valeur de ω (et donc de la fréquence).
- Dans le cas d'une modulation de fréquence, la pulsation ω est une fonction linéaire du signal modulant. Soit $\psi = \omega \cdot t + \varphi$.

Par définition, la fréquence est égale à la dérivée de la phase instantanée :

$$\frac{d\psi}{dt} = \frac{d(\omega \cdot t + \varphi)}{dt} = \frac{d(2 \cdot \pi \cdot f \cdot t + \varphi)}{dt} = 2 \cdot \pi \cdot f \Rightarrow f = \frac{1}{2 \cdot \pi} \cdot \frac{d\psi}{dt}$$

Le fréquence f désigne une quantité variable que nous notons par la suite $f(t)$.

$$\psi(t) = \omega \cdot t + \frac{\Delta f}{f_1} \sin(\omega_1 \cdot t) + \Phi$$



Prenons $m(t) = B \cdot \cos(2 \cdot \pi \cdot f_1 \cdot t + \varphi_1) = B \cdot \cos(\omega_1 \cdot t + \varphi_1)$

$$n(t) = A \cdot \sin\left(\omega \cdot t + \frac{\Delta f}{f_1} \sin(\omega_1 \cdot t) + \Phi\right)$$

L'**indice de modulation en fréquence** s'écrit :

$$m_F = \frac{\Delta f}{f_1}$$

L'excursion en fréquence est $\Delta f = m_F \cdot f_1$

Pour les besoins de l'étude spectrale, développons l'expression d'un signal modulé en fréquence par une sinusoïde :

$$\begin{aligned} n(t) &= A \cdot \sin(\omega \cdot t + \Phi + m_F \cdot \sin(\omega_1 \cdot t)) \\ &= A \cdot [\sin(\omega \cdot t + \Phi) \cdot \cos(m_F \cdot \sin(\omega_1 \cdot t)) + \sin(m_F \cdot \sin(\omega_1 \cdot t)) \cdot \cos(\omega \cdot t + \Phi)] \end{aligned}$$

En utilisant les fonctions de Bessel de première espèce $J_n(x)$, nous pouvons réécrire cette relation sous la forme :

$$n(t) = A \cdot \left[J_0(m_F) \cdot \sin(\omega \cdot t + \Phi) + \sum_{k=1}^{\infty} J_k(m_F) \cdot \sin(\omega \cdot t + n \cdot \omega_1 \cdot t) + (-1)^k \cdot \sin(\omega \cdot t - n \cdot \omega_1 \cdot t) \right]$$

Le spectre contient donc une infinité de raies dont les amplitudes répondent à des relations incluant des fonctions de Bessel.

Dans le cas d'une **modulation de phase**, la phase φ de la porteuse est une fonction linéaire de celle du signal modulant. Soit $\Delta\varphi = \Delta\vartheta \cdot m(t) = \Delta\vartheta \cdot \cos(\omega_1 \cdot t)$.

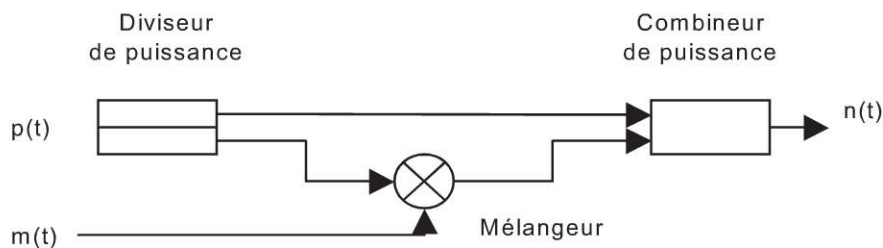
Le signal modulé en phase résultant est de la forme suivante :

$$n(t) = A \cdot \sin(\omega \cdot t + \Delta\vartheta \cdot \cos(\omega_1 \cdot t) + \varphi)$$

Similairement à la modulation de fréquence, la modulation de phase a un spectre constitué d'une infinité de composantes spectrales à des pulsations $\omega \pm n \cdot \omega_1$ et dont l'amplitude se calcule avec des relations incluant des fonctions de Bessel de première espèce.

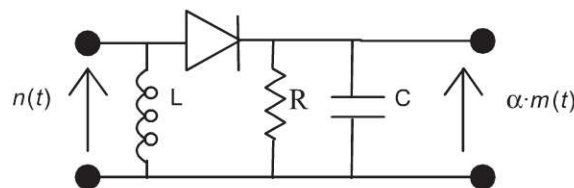
3. EN PRATIQUE

a) Schéma synoptique d'un modulateur d'amplitude



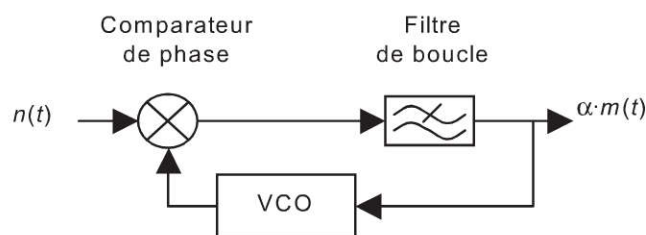
b) Schéma électrique d'un démodulateur d'amplitude

La démodulation d'amplitude repose sur la détection d'enveloppe : la diode effectue un redressement simple alternance (fiches 31 et 72) et les composants R et C forment un filtre passe-bas filtrant la composante fréquentielle de l'enveloppe. Le self sert à polariser la diode.



c) Schéma électrique d'un démodulateur FM

Le montage à PLL (fiche 64) génère une tension de basse fréquence proportionnelle à la phase de signal FM reçu : ce signal correspond à l'information transportée car ce système convertit la phase en niveau de tension.



67 Modulations numériques

Mots clés

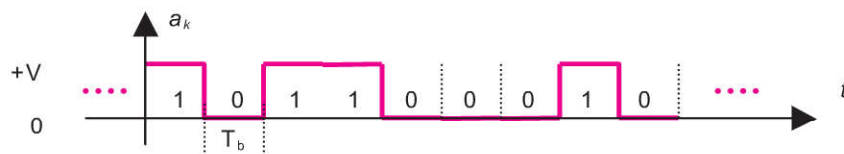
Modulations OOK, FSK, CPFSK, MSK, BPSK, QPSK, QAM-N.

1. EN QUELQUES MOTS

Les modulations numériques désignent les techniques de transmission d'informations numériques sur une fréquence porteuse. Tout comme pour les modulations analogiques (fiche 66), les modulations numériques modulent la porteuse $p(t) = A \cdot \cos(\omega \cdot t + \phi)$ en amplitude, en fréquence ou en phase (et éventuellement simultanément en amplitude et en phase).

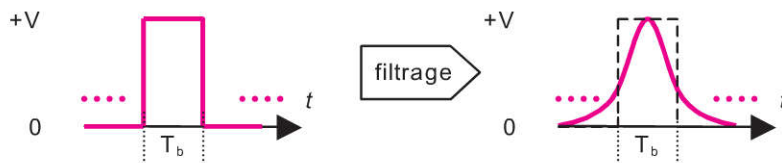
2. CE QU'IL FAUT RETENIR

Nous supposons dans un premier temps que l'information numérique à transmettre est sous forme sérialisée conformément au schéma suivant :



Le débit binaire est l'inverse de la durée d'un bit : $D = \frac{1}{T_b}$

Le débit binaire s'exprime en bits par seconde ou en bauds qui est une unité équivalente. Le filtrage visant à limiter l'occupation spectrale du signal modulé modifie l'allure de chaque impulsion :

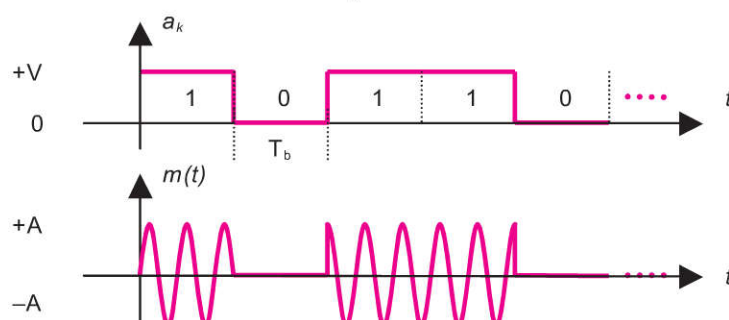


Toutes les impulsions s'étalant dans le temps, il peut en résulter un phénomène d'interférences inter-symboles (ISI), c'est-à-dire que les impulsions filtrées sont trop proches les une des autres pour distinguer un 0 d'un 1 logique.

a) Modulation d'amplitude tout-ou-rien (OOK, On Off Keying)

La modulation OOK présente des similitudes avec la transmission du code Morse. Elle correspond à la multiplication directe du train binaire a_k par la porteuse, il en résulte un signal modulé

$$n(t) = a_k \cdot A \cdot \sin(\omega \cdot t + \phi)$$



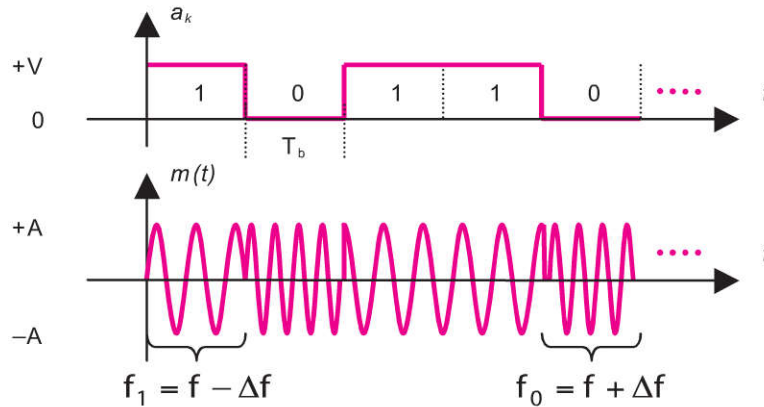
b) Modulation de fréquence FSK (Frequency Shift Keying)

La modulation FSK associe une fréquence de porteuse f_0 pour un bit à 0 et une fréquence différente f_1 pour un bit à 1. Soient f la moyenne de ces deux fréquences et Δf leur différence par rapport à f .

$$f_1 = f - \Delta f \Leftrightarrow \omega_1 = \omega - \Delta\omega$$

$$f_0 = f + \Delta f \Leftrightarrow \omega_0 = \omega + \Delta\omega$$

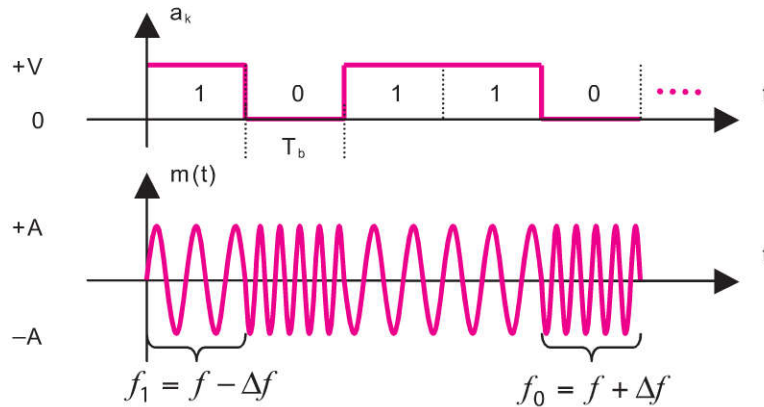
Le signal modulé peut s'écrire $n(t) = A \cdot \sin([\omega + (2 \cdot a_k - 1) \cdot \Delta\omega] \cdot t)$



Les sinusoïdes de fréquences f_0 et f_1 sont issues de deux générateurs distincts. C'est pourquoi un changement de valeur de bit produit un saut de phase : le signal modulé présente des discontinuités. Cela élargit l'encombrement spectral : chaque discontinuité produit des fréquences élevées si le signal résultant est analysé dans le domaine fréquentiel.

c) Modulation de fréquence CPFSK (Continuous Phase FSK)

La modulation CPFSK palie au problème de discontinuités de phase.



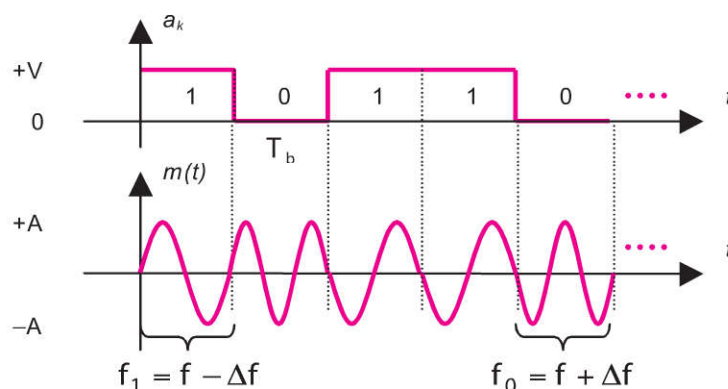
Dans ce cas, à débit égal, l'encombrement spectral est plus réduit que dans celui de la FSK simple, l'efficacité spectrale est donc meilleure.

d) Modulation de fréquence MSK (Minimum Shift Keying)

La modulation MSK correspond à une configuration particulière de la CPFSK, celle où :

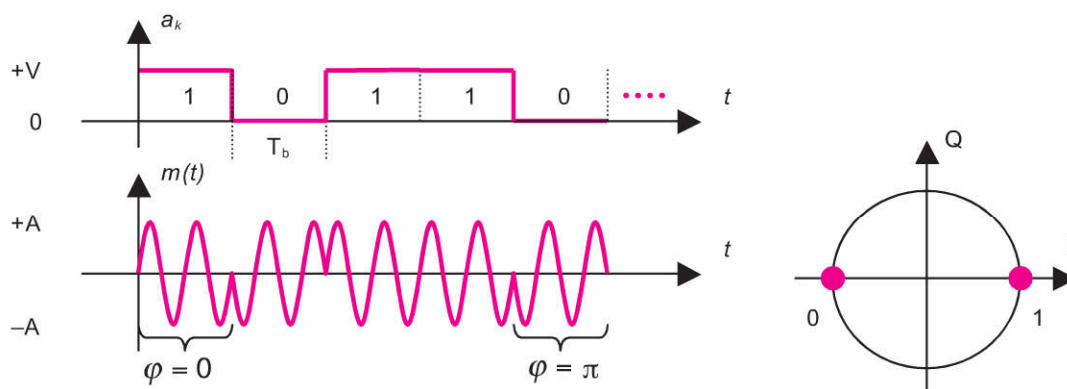
$$\frac{2 \cdot \Delta f}{D} = \frac{f_0 - f_1}{D} = 0,5$$

Par conséquent, la différence pour coder un 0 ou un 1 est d'exactement une demi-période de la porteuse comme nous pouvons le voir sur le diagramme suivant :



e) Modulation de phase 2-PSK ou BPSK (Binary Phase Shift Keying)

La modulation BPSK associe un état de phase à chaque valeur binaire à transmettre. Dans le cas présenté, 1 correspond à une phase nulle et 0 à une phase de π .

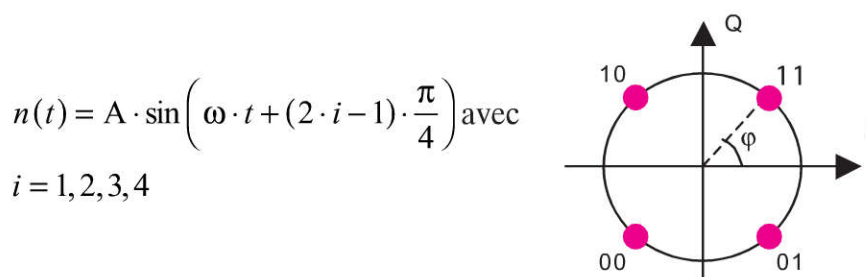


La répartition des points indiquant la correspondance entre la valeur binaire et la phase dans le plan complexe est appelée **constellation**.

Le signal modulé BPSK répond à l'équation : $m(t) = A \cdot \sin(\omega \cdot t + (1 - a_k) \cdot \pi)$

f) Modulation de phase 4-PSK ou QPSK (Quadrature Phase Shift Keying)

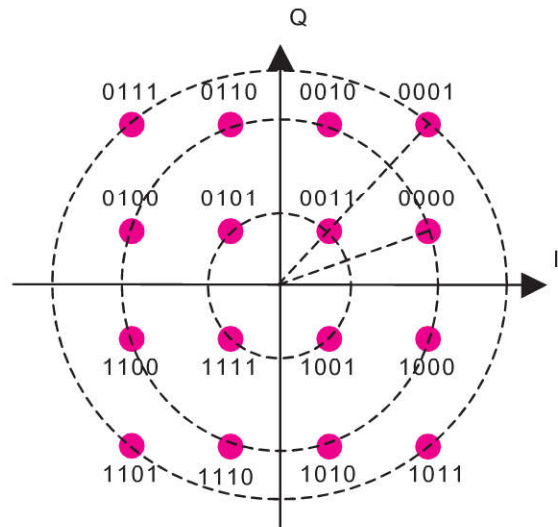
La modulation BPSK affecte un état de phase par valeur binaire : deux états de phase suffisent pour le transport de l'information. En augmentant le nombre d'états de phase permis, le débit correspondant est amélioré. Au même titre que pour la modulation BPSK, seule la phase du signal modulé est modifiée en fonction du mot binaire à transmettre. La modulation QPSK double le débit par rapport à la modulation BPSK : à chacun des quatre états de phase correspond un mot binaire de 2 bits. L'information numérique à transmettre n'est plus sérialisée mais transmise par paquets de 2 bits :



i	φ	a_k
1	$\pi / 4$	11
2	$3\pi / 4$	10
3	$5\pi / 4$	00
4	$7\pi / 4$	01

g) Modulations QAM (Quadrature Amplitude Modulation)

Afin d'augmenter le débit de la modulation QPSK, il est possible d'augmenter le nombre d'états de phase : on ajoute alors un bit à chaque doublement du nombre de phases. Ces modulation s'appellent N-PSK, avec N le nombre d'états de phase. Pour augmenter encore le débit tout en évitant de réduire l'écart entre les phases permises (ce qui aboutirait rapidement à l'apparition d'ambiguïtés), les modulations QAM associent les principes des modulations de phase et d'amplitude. Le diagramme suivant illustre la répartition des phases et des amplitudes par rapport aux 16 mots de 4 bits d'une 16-QAM :



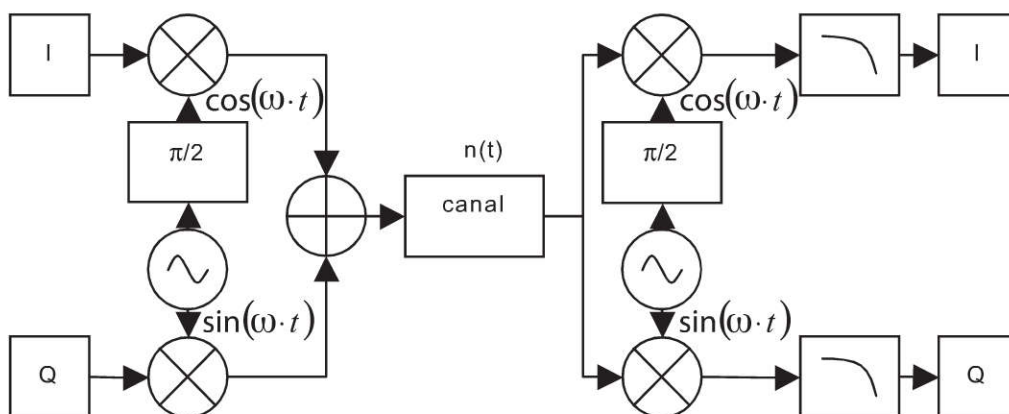
Différentes formes de modulations N-QAM sont possibles : la répartition présentée est un cas particulier où les points de la constellation sont disposés en carré, mais la forme pourrait être rectangulaire par exemple.

3. EN PRATIQUE

La **modulation** OOK peut tout simplement être mise en place en actionnant un interrupteur avec le train d'informations binaires. En fonction du bit à transmettre, l'interrupteur, ouvert ou fermé, laisse passer le signal d'un oscillateur réglé sur la fréquence porteuse. La **démodulation** met en œuvre une détection d'enveloppe : le signal reçu est filtré par un passe-bas afin d'éliminer la composante correspondant à la porteuse.

La modulation FSK simple consiste à commuter entre deux sources de fréquences différentes. Pour obtenir une phase continue (CPFSK), nous n'utilisons pas deux générateurs, mais un seul. Ce générateur unique est un VCO (*Voltage Controlled Oscillator*, [fiche 64](#)). La démodulation suppose de multiplier le signal dans deux branches différentes contenant chacune un oscillateur réglé sur les deux fréquences contenues dans le signal.

Les modulations de phase, et particulièrement les modulations QPSK et N-QAM, font appel à des modulateurs et démodulateurs dits « I et Q ». La composante réelle (I pour In-phase) et la composante imaginaire (Q pour Quadrature) sont générées séparément, modulées par deux sinusoïdes en quadrature puis additionnées avant d'être émises. La démarche réciproque est appliquée pour la réception (démodulation).



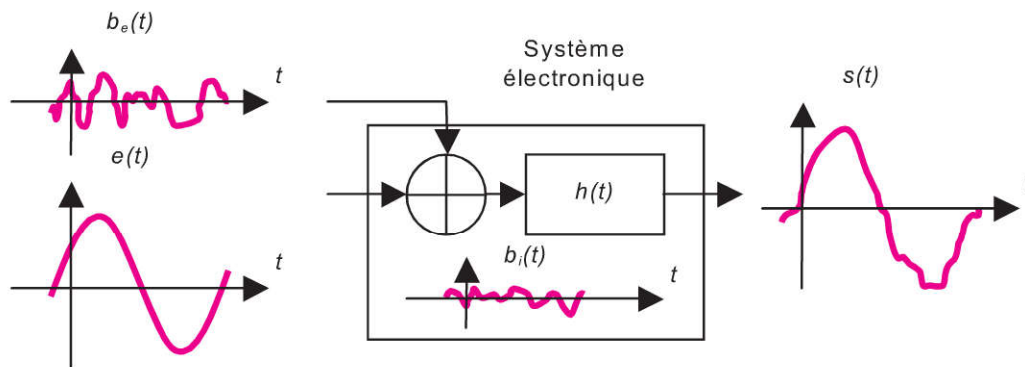
68 Le bruit

Mots clés

Bruit blanc, bruit rose, bruit de grenaille, bruit thermique, mouvement brownien, SNR.

1. EN QUELQUES MOTS

Le signal est une manifestation physique (typiquement dans notre cas, un courant électrique) qui transporte une information. Le bruit se superpose au signal et n'apporte pas d'information, bien au contraire : il perturbe la bonne « lecture » du signal utile. Le bruit peut provenir de l'extérieur du système (perturbations électromagnétiques dues au réseau d'alimentation électrique 50 Hz, bruit de fond cosmique), soit de l'intérieur du système, c'est-à-dire des composants eux-mêmes.



2. CE QU'IL FAUT RETENIR

a) Densité spectrale de puissance

La densité spectrale de puissance correspond à la répartition fréquentielle de la puissance de bruit. Elle est égale à la transformée de Fourier de la fonction d'autocorrélation du bruit. Elle s'exprime en V^2/Hz ou en A^2/Hz . Par analogie avec la lumière, on parle de bruit blanc lorsque la densité spectrale de puissance (DSP) est constante dans une bande de fréquence considérée. Le bruit rose (ou bruit en « 1 sur f ») n'a pas une DSP constante, mais proportionnelle à $1/f$: il tend vers 0 pour les hautes fréquences et vers l'infini pour les fréquences les plus basses.

b) Les sources de bruit

► Bruit thermique

Les électrons peuvent être soumis à un mouvement d'ensemble constituant le courant électrique. Individuellement, ces électrons présentent également un mouvement aléatoire qui induit des variations de potentiel assimilables à du bruit. À titre d'exemple, le bruit thermique (ou bruit de Johnson-Nyquist) $X(f)$ dans une résistance est un bruit blanc dont la densité spectrale de puissance (constante) est sensiblement égal à :

$$X_{th}(f)_{A^2 \cdot Hz^{-1}} \approx 4 \cdot k \cdot T \cdot R$$

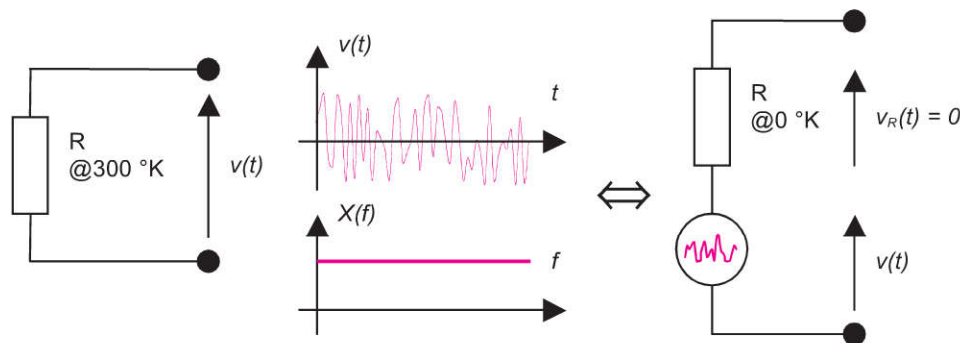
Avec

k : la constante de Boltzmann soit $1,38 \cdot 10^{-23} \text{ J} \cdot \text{K}^{-1}$

T : la température en degrés Kelvin

R : la résistance en ohms

H : la constante de Planck soit $6,6 \cdot 10^{-34} \text{ J} \cdot \text{s}$



► Bruit de grenaille

Dans un semi-conducteur traversé par un courant électrique, il existe un phénomène de générations/recombinaisons aléatoire de porteurs au niveau d'une jonction. Le bruit généré, de type bruit blanc, est proportionnel au courant qui circule dans le composant.

$$X_g(f)_{A^2 \cdot Hz^{-1}} = 2 \cdot q \cdot \overline{i(t)}$$

$\overline{i(t)}$ est le courant moyen traversant la jonction, en ampères.

q est la charge élémentaire.

► Bruit de scintillation

Ce bruit (également appelé *flicker* en anglais) est essentiellement dû à des impuretés dans les semi-conducteurs provoquant des recombinaisons électron-trou aléatoires. Ce phénomène provient du courant continu dans les composants (polarisation). C'est un bruit rose pour lequel la fréquence de coude (fréquence pour laquelle le bruit blanc est de même puissance que le bruit de scintillation) dépend de la technologie du composant. Plus précisément, ce bruit souvent assimilé à un bruit en $1/f$ répond en fait plutôt à la forme suivante :

$$X_s(f)_{A^2 \cdot Hz^{-1}} = C/f^\alpha \text{ avec } C \text{ une constante et } 0,7 \leq \alpha \leq 1,3$$

► Bruit d'avalanche

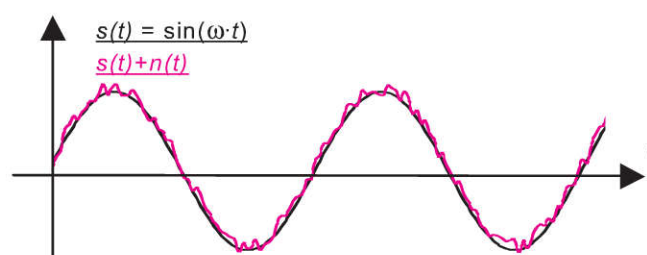
Dans les semi-conducteurs, le champ électrique est tel que des électrons peuvent faire « sauter » des électrons de valence dans la bande de conduction, qui deviennent autant d'électrons libres supplémentaires. Plus le champ électrique est important, plus les électrons sont globalement accélérés, plus ce phénomène prend de l'importance, particulièrement à proximité de la zone d'avalanche des jonctions PN (fiche 30).

3. EN PRATIQUE

La quantité de bruit est évaluée en indiquant le rapport signal à bruit ou SNR (*Signal to Noise Ratio*). Cette grandeur est exprimée en dB :

$$SNR_{dB} = 10 \cdot \log \left(\frac{P_s}{P_b} \right)$$

Une valeur de SNR positive élevée indique que le bruit influence peu le signal. Une valeur négative signifie que la puissance du bruit est plus forte que celle du signal.



69 Introduction à l'électronique de puissance

Mots clés

Source, interrupteur, énergie, transformation, puissance

1. EN QUELQUES MOTS

L'électronique de puissance est l'une des branches de l'électrotechnique. Elle traite l'énergie électrique par voie conversion statique (fiche 72). C'est aujourd'hui un domaine en plein essor grâce aux nouveaux champs d'applications liés au développement durable et aux énergies renouvelables.

2. CE QU'IL FAUT RETENIR

a) Applications

Les applications de l'électronique de puissance sont diverses et concernent aussi bien les particuliers que les utilisations professionnelles.

► Applications domestiques

- Alimentations à découpage dans les téléviseurs, les ordinateurs...
- Chargeurs des batteries pour les téléphones et les ordinateurs portables
- Dispositifs qui font varier la vitesse de rotation des machines à laver, aspirateurs, mixers, etc.

► Applications industrielles

- Alimentation des moteurs pour variation de vitesse.
- Alimentation de secours et alimentation sans interruption pour les communications, la production pétrolière, etc.
- Alimentation des lampes fluorescentes basse consommation.
- Alimentation pour laser, radar...
- Générateurs d'ultrasons ou d'électricité utilisés dans le domaine médical.

On utilise l'électronique de puissance pour la fabrication des dispositifs terrestres dans le transport (tramway, TGV, voiture hybride) de même qu'on peut l'utiliser dans les domaines naval et aéronautique.

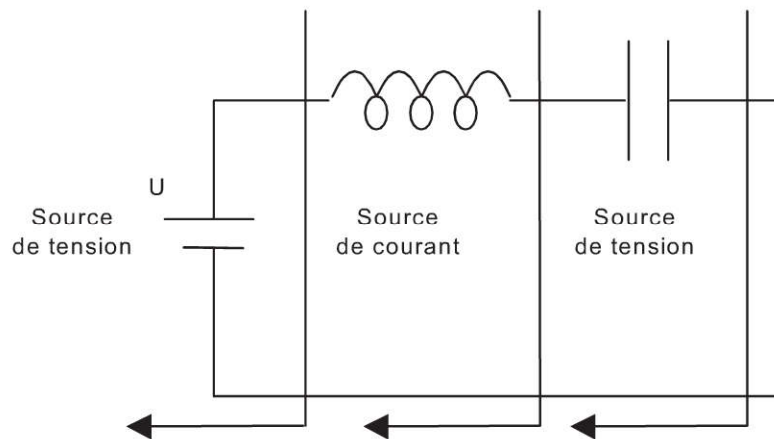
3. EN PRATIQUE

a) Source électrique

On peut distinguer deux types de source : les sources continues et les sources alternatives pouvant chacune être source de tension ou de courant. Comment distinguer une source de tension d'une source de courant ?

Si la source est en série avec une inductance, c'est une source de courant.

Une source de tension, quant à elle, est en parallèle avec un condensateur.



Nous pouvons utiliser le réseau d'énergie (EDF), comme une tension alternative (source alternative en France $220\text{ V}_{\text{eff}}/50\text{ Hz}$). Ces valeurs ne sont pas pareilles pour tous les pays, par exemple : aux États-Unis nous avons 120 V et 60 Hz ou au Japon 100 V et 50 Hz , etc.

La deuxième source alternative peut être fournie depuis un onduleur, qui nous fournit une tension alternative en sortie. Nous discuterons de ce dispositif dans les fiches suivantes. Pour avoir des sources continues, nous avons plusieurs choix : les condensateurs, les batteries et les panneaux solaires, qui nous fournissent une source continue et la valeur de cette source est variable selon notre besoin. Pour charger les batteries, nous pouvons également utiliser les éoliennes. L'énergie éolienne utilise la force du vent pour faire tourner des aérogénérateurs pour fournir des énergies comme une source. Actuellement, il y a 97 GW d'énergie éolienne installés dans le monde.

b) Réversibilité des sources

Une source est réversible en courant si le courant qui la traverse peut s'inverser. Une source est appelée réversible en tension si la tension à ses bornes peut changer de signe. Par exemple, une batterie est une source réversible en courant (charge et décharge) et non réversible en tension.

70 Modèles simplifiés des semi-conducteurs de puissance

Mots clés

Interrupteur, semi-conducteur, schéma équivalent.

1. EN QUELQUES MOTS

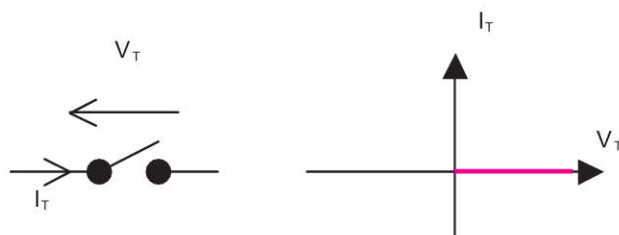
Le modèle de l'interrupteur permet de simplifier considérablement l'analyse de circuits mettant en œuvre des composants à base de semi-conducteurs, à l'image de l'étude présentée dans la [fiche 30](#). Les interrupteurs à semi-conducteurs statiques (c'est-à-dire sans mouvement mécanique, [fiches 70](#) et [72](#)) permettent d'ouvrir et de fermer un circuit alimentant une charge électrique.

2. CE QU'IL FAUT RETENIR

Un convertisseur statique nécessite des interrupteurs statiques. Ils se trouvent dans l'un des deux états possibles :

- état passant ou fermeture ;
- état bloqué ou ouverture.

On peut symboliser ce fonctionnement à l'aide d'un modèle de type « fonction de transfert ».



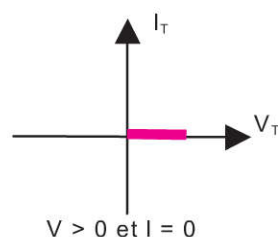
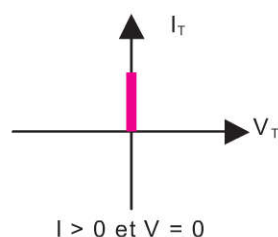
Il existe quatre configurations pour les interrupteurs statiques.

- **Un segment** : une tension appliquée aux bornes de l'interrupteur sans passage de courant est équivalente à un circuit ouvert et réciproquement, un courant circulant avec une tension nulle est un court-circuit (fil).
- **Deux segments** : interrupteur monodirectionnel en courant et en tension.
- **Trois segments** : interrupteur monodirectionnel en courant ou en tension et bidirectionnel en tension ou en courant.
- **Quatre segments** : interrupteur bidirectionnel en courant et en tension.

a) Caractéristique d'un interrupteur à un segment :

Un interrupteur monosegment peut être considéré comme un court-circuit ou comme un circuit ouvert :

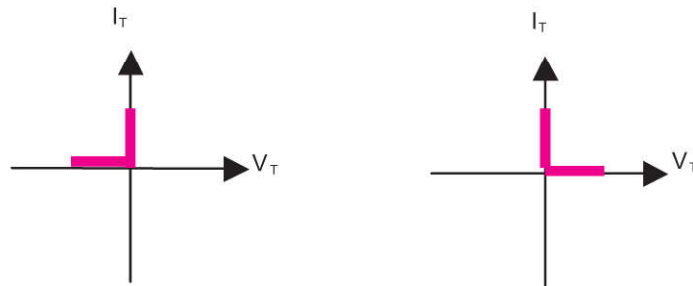
- $V_T > 0$ et $I_T = 0$
- $I_T > 0$ et $V_T = 0$



b) Caractéristique d'un interrupteur à deux segments

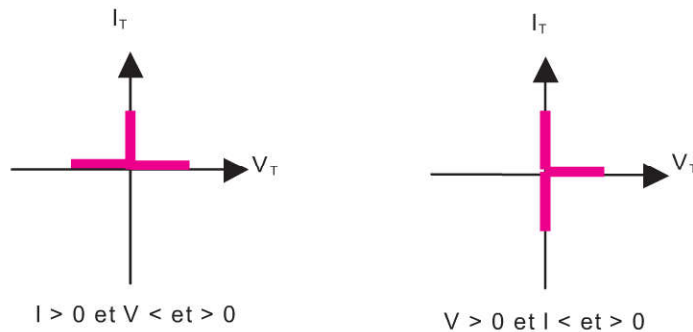
Les diodes, les transistors bipolaires, les IGBT et les MOSFET entrent dans cette catégorie.

- $V_T > 0$ et $I_T < 0$
- $I_T > 0$ et $V_T > 0$

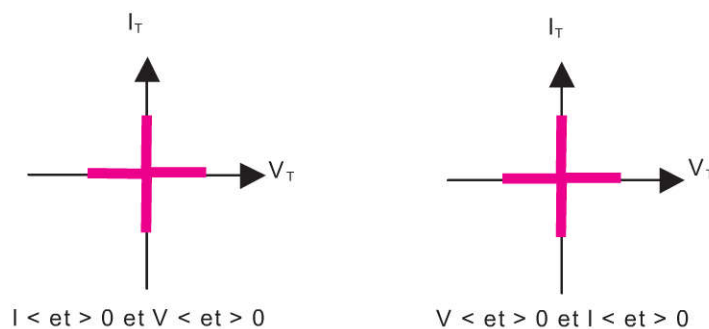
**c) Caractéristique d'un interrupteur à trois segments**

Cette catégorie regroupe les Thyristors et les IGBTs –diode.

- $V_T > 0$ et $I_T > 0$ et < 0
- $I_T > 0$ et $V_T > 0$ et < 0

**d) Caractéristique d'un interrupteur à quatre segments**

Les Triacs et tous les interrupteurs bidirectionnelles sont dans cette catégorie : $V_T > 0$ et < 0 et $I_T > 0$ et < 0 .

**3. EN PRATIQUE**

Ce formalisme, très simple, voire simpliste, permet d'aborder avec un certain recul les aspects techniques présentés dans les [fiches 70](#) et [72](#).

71 Composants semi-conducteurs de puissance

Mots clés

Transistor bipolaire, thyristor, GTO, TRIAC, MOSFET, IGBT, IGCT

1. EN QUELQUES MOTS

Les semi-conducteurs sont indispensables pour réaliser des fonctions électroniques de puissance : dans les interrupteurs statiques, on trouve un assemblage de jonctions PN. Rappelons que la zone N contient des ions positifs fixes et des électrons mobiles (porteurs) et que la zone P contient des ions négatifs fixes et des trous mobiles.

2. CE QU'IL FAUT RETENIR

a) Transistor bipolaire

Le transistor bipolaire est un composant électronique actif qui peut être utilisé comme interrupteur (mode de commutation). Le transistor amplifie un courant : il s'utilise pour l'ouverture ou la fermeture des circuits, d'oscillateurs ou pour moduler un signal dans les oscillateurs entre autres applications. Au sein d'un transistor, le courant passant entre le collecteur et l'émetteur est contrôlé par le courant appliqué à la base. Les transistors sont utilisés à peu près partout comme, par exemple, dans les microprocesseurs : il y en a 291 millions dans un Core 2 Duo.

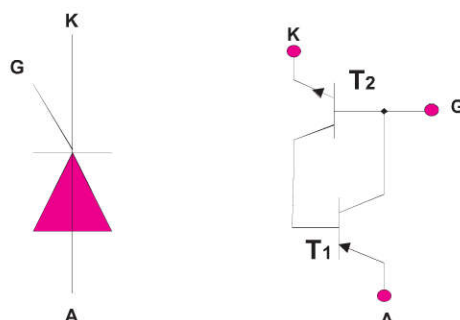
► Thyristor

Le thyristor est un interrupteur électronique semi-conducteur qui se compose de quatre couches de silicium P-N-P-N. Il peut être modélisé par deux transistors PNP et NPN. Il ne laisse passer le courant que dans un seul sens, comme la diode, et peut être commandé à l'allumage et pas à l'extinction, mais il peut être bloqué. Nous avons deux états : amorçage et blocage. Si la tension anode-cathode est positive $V_{AK} > 0$ et le courant de gâchette (de G vers K) $I_G > I_{G\max}$, dans ce cas le thyristor est dans l'état d'amorçage. En revanche, dans un thyristor, nous avons deux types de blocage :

- blocage naturel par annulation, le courant I_{AK} (anode-cathode) comme un pont mixte ou un pont complet ; dans ce cas, le thyristor fonctionne en courant redressé et son blocage est naturel à chaque période ;
- blocage forcé par inversion V_{AK} , par exemple quand il est au mode de fonctionnement en courant continu (fiche 72).

Remarque importante

Une inductance placée en série avec le thyristor ralentit la vitesse de croissance du courant principal au début de l'amorçage du thyristor.



Dans un thyristor, I_G est le courant de la gâchette (positif lorsqu'il rentre), V_{AK} la tension entre l'anode et la cathode du thyristor et I_{AK} le courant considéré positif lorsqu'il traverse le thyristor de l'anode vers la cathode.

Ci-dessous, on présente les équations nécessaires pour un thyristor :

$$I_{C1} = \beta_1 \cdot I_{B1} + I'_{C1}$$

$$I_{C2} = \beta_2 \cdot I_{B2} + I'_{C2}$$

Ensuite, on trouve le courant d'anode :

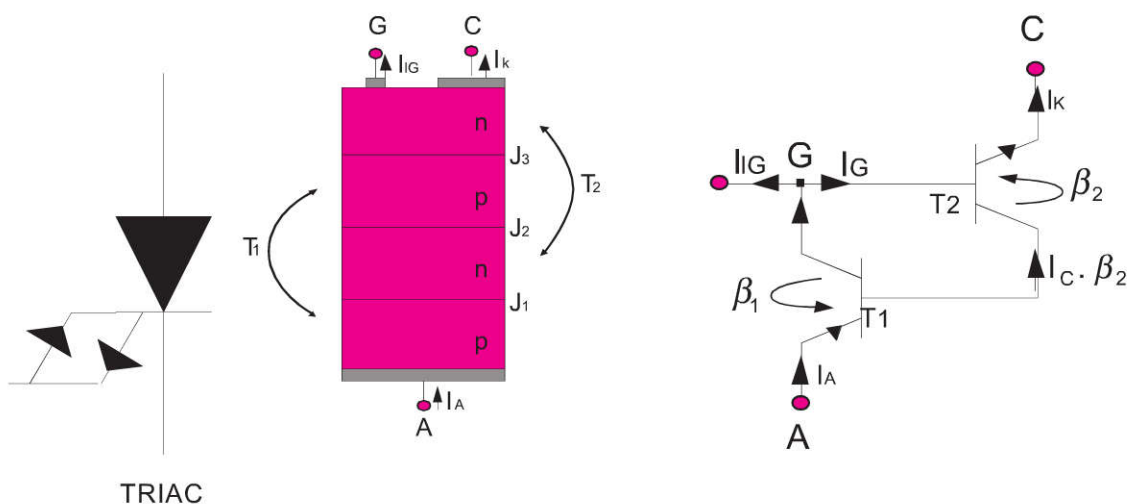
$$I_A = I_{B2} + I_{C2} = \frac{(1 + \beta_2) \cdot \beta_1 \cdot I'_{C2} + (1 + \beta_2) \cdot I'_{C1}}{1 - \beta_1 \cdot \beta_2} + I'_{C1}$$

I_A est le courant d'anode et I_{B1} et I_{B2} sont respectivement le courant de la base du premier transistor et le courant de base du deuxième transistor. I_{C1} et I_{C2} sont respectivement le courant de premier collecteur et le courant de deuxième collecteur. I'_{C1} et I'_{C2} sont les courants de fuite des deux transistors, β_1 et β_2 sont les gains des deux transistors. Dans l'équation ci-dessus, si on suppose $I_G = 0$, en conséquence I_{C1} et β_1 sont faibles.

Ces composants sont utilisés pour la conversion alternatif-continu (fiche 72), les variateurs de vitesse des petits moteurs électriques, etc.

► GTO

Les *Gate Turn-Off Thyristor* (GTO, thyristor blocable par la gâchette) sont des thyristors à extinction. Ce sont des interrupteurs électroniques utilisés dans les dispositifs de forte puissance adaptés aux courants élevés et aux grandes vitesses de commutation. On les utilise dans les onduleurs et les hacheurs (fiche 72). Le GTO est structuellement identique à un thyristor, il se compose de trois électrodes (anode A, cathode K et électrode de gâchette G pour la commande) et se compose de quatre couches P, N, P, N.



► TRIAC

Le Triac (*triode for alternating current*) est un composant électronique équivalent à la mise en parallèle de deux thyristors montés tête-bêche. Un triac est, contrairement au thyristor, un composant bidirectionnel, qui peut laisser passer le courant dans les deux sens. C'est un interrupteur commandé et sa gâchette est la branche en diagonale. Le Triac se compose de deux structures P1 N1

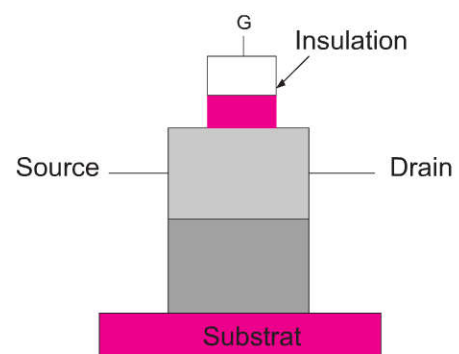
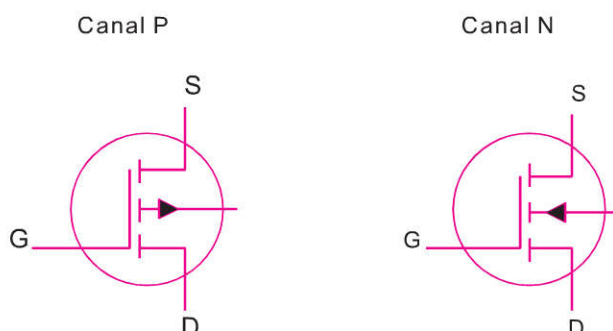
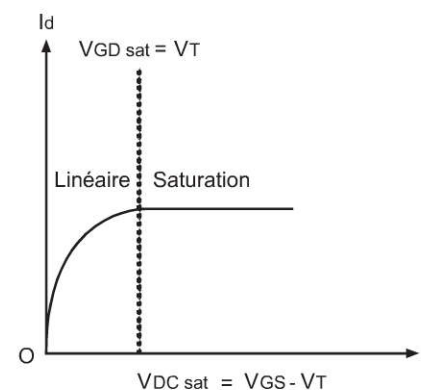
P2 N2 de thyristor et P2 N1 P1 N4, qui sont montés en parallèle inverse. A1 est également reliée à P2 et A2 reliée également à une couche supplémentaire N4. G est également reliée à une couche supplémentaire N3. Ci-dessous, on montre le schéma symbolique et la structure réelle d'un triac.



► MOSFET

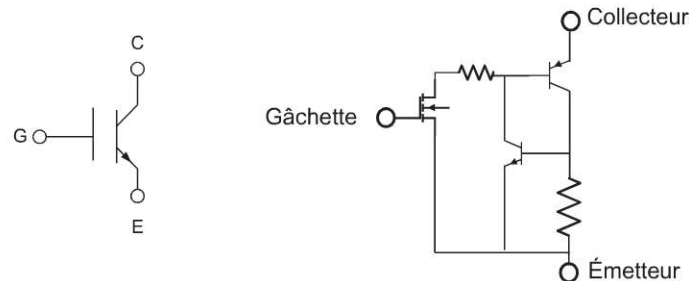
Les MOSFET (*Metal Oxide Semiconductor Field Effect Transistor*, transistor à effet de champ à métal oxyde semiconducteur) ont leur grille isolée du canal par une couche de dioxyde de silicium (SiO_2). Les transistors à effet de champ (fiche 42) de type MOS se composent de quatre électrodes : la source (S), point de départ des porteurs, le drain (D) point de collecte des porteurs, la grille (G) et le substrat (B). Nous avons deux types de MOSFET : à canal N et à canal P. L'intensité du courant circule toujours entre la source et le drain, et les tensions sont mesurées par rapport à la source. Dans un MOSFET, nous avons trois zones de fonctionnement :

- **Zone bloquée ou de coupure** lorsque $V_{GS} < V_T$, $V_{GD} < V_T$ et $V_{DS} > 0$, $I_{DS} = 0$. Le composant se comporte comme un interrupteur ouvert. V_T est la tension de seuil et V_{GS} la tension entre la grille et la source.
- **Zone linéaire ou triode** : $V_{GS} - V_T > V_{DS}$, $V_{GS} \geq V_T$ et $V_{GD} > V_T$ avec $V_{DS} > 0$
- **Zone de saturation** : dans cette zone nous avons $V_{GS} \geq V_T$, $V_{GD} < V_T$ et $V_{DS} > 0$. I_{DS} est indépendante de V_{DS} et le MOSFET se comporte comme un générateur de courant commandé en tension. Le courant entre le drain et la source peut être calculé selon la loi $I_{DS} = (K/2) \cdot (V_{GS} - V_T)^2$ (K est une constante caractéristique du composant).



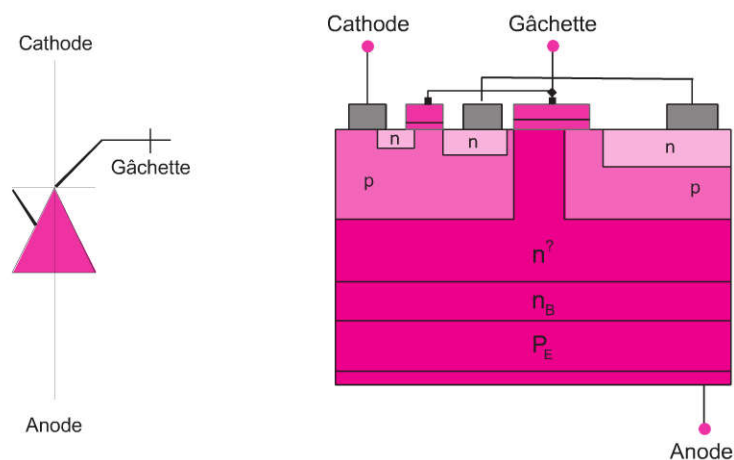
► IGBT

L'IGBT (*Insulated Gate Bipolar Transistor*) est un transistor bipolaire à porte isolée. Il se compose d'un transistor bipolaire et d'un MOSFET. Sa faiblesse par rapport au MOSFET est la vitesse, notamment lors du passage de l'état passant à l'état bloqué ce qui limite sa fréquence de commutation à quelques dizaines de kHz. Les applications usuelles incluent la conversion statique (fiche 72), ou encore l'alimentation des véhicules électriques. Si $V_{GE} > V_T$ (tension de seuil), le composant est passant, il est bloqué dans le cas contraire.



► IGCT

Le thyristor commuté par porte intégré (*Integrated-Gate commutated Thyristor*) est un nouveau dispositif de haute puissance utilisé dans les équipements industriels. Un IGCT est une évolution de la famille des thyristors des GTO. Offrant plus de contrôles, il a des paramètres dynamiques très rapides pour passer du mode activé à la désactivation. Il présente une fréquence de commutation élevée, de faibles pertes de commutation et ne nécessite pas de circuits de protection. La diode et la gâchette sont intégrées pour diminuer l'encombrement et permettre la construction d'appareils modulaires et compacts.



3. EN PRATIQUE

On peut résumer ce chapitre en montrant la table de comparaison des interrupteurs statiques.

Composant	Commande	Fréquence Max	Blocage	Pertes en conduction	Pertes en commutation
Diode	Non	élevée	> 10 kV	faibles	nulles
Bipolaire	On/Off	10 kHz	1,2 kV	faibles	élevées
Thyristor	On/Off	< 1 kHz	> 10 kV	faibles	élevées
GTO	On/Off	< 1 kHz	> 10 kV	faibles	élevées
MOSFET	On/Off	250 kHz	600 V	élevées	faibles
IGBT	On/Off	25 kHz	4,5 kV	moyennes	moyennes
IGCT	On/Off	40 kHz	4,5 kV	faibles	très faibles

72 Convertisseurs statiques

Mots clés

Transformation, source, énergie, gradateurs, cycloconvertisseurs, redresseurs, hacheurs, onduleurs, pont de diodes, rapport cyclique

1. EN QUELQUES MOTS

Tout dispositif qui modifie les caractéristiques d'une source d'énergie d'entrée s'appelle un convertisseur statique. Les convertisseurs statiques permettent d'adapter la source d'énergie électrique à un récepteur donné. Ce sont des circuits électriques utilisant des interrupteurs statiques (semi-conducteurs) en régime de découpage pour traiter l'énergie électrique à haut rendement et assurer l'alimentation d'une charge électrique. Il existe quatre types de convertisseurs statiques :

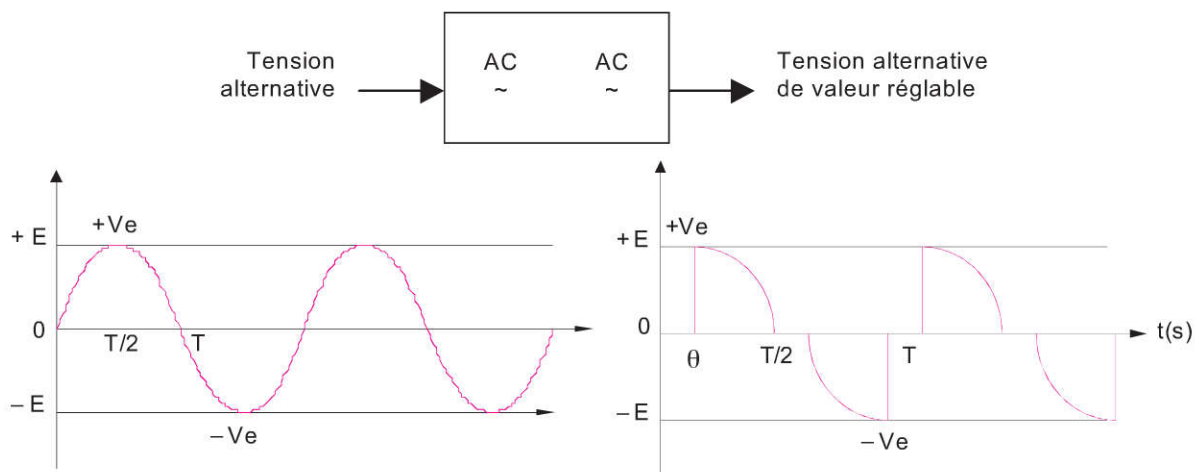
- AC-AC : gradateurs ou cycloconvertisseurs ;
- AC-DC : redresseurs ;
- DC-DC : hacheurs ;
- DC-AC : onduleurs.

2. CE QU'IL FAUT RETENIR

a) Convertisseur AC-AC

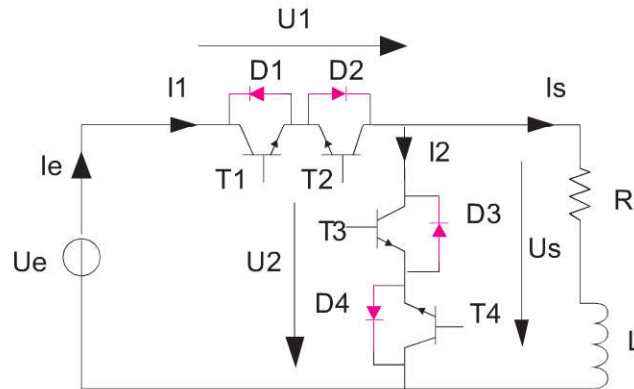
Ce convertisseur peut mettre en relation une source de tension alternative en entrée à une source de courant alternatif en sortie. Nous pouvons distinguer deux catégories de convertisseurs statiques AC-AC :

- les gradateurs à fréquence invariable ou fixe ;
- les cycloconvertisseurs à fréquence variable.



Un convertisseur AC-AC peut être commandé à la fermeture et à l'ouverture ou à la fermeture seulement. On peut réaliser un dispositif AC-AC à la fermeture à l'aide de thyristors. Par contre, pour un dispositif commandé à la fermeture et à l'ouverture, on doit utiliser des IGBT associés avec des diodes. Ci-dessous est illustré un convertisseur AC/AC très courant commandé à la fermeture et à l'ouverture avec une fréquence invariable. Il est constitué d'interrupteurs bidirectionnels. Chaque interrupteur comprend un IGBT avec une diode antiparallèle qui est à l'intérieur de

l'IGBT. Les interrupteurs dans ce schéma sont deux IGBT bidirectionnels. U_e est la tension du réseau qui est telle que $U_e = \hat{U} \cdot \sin(\omega \cdot t)$.

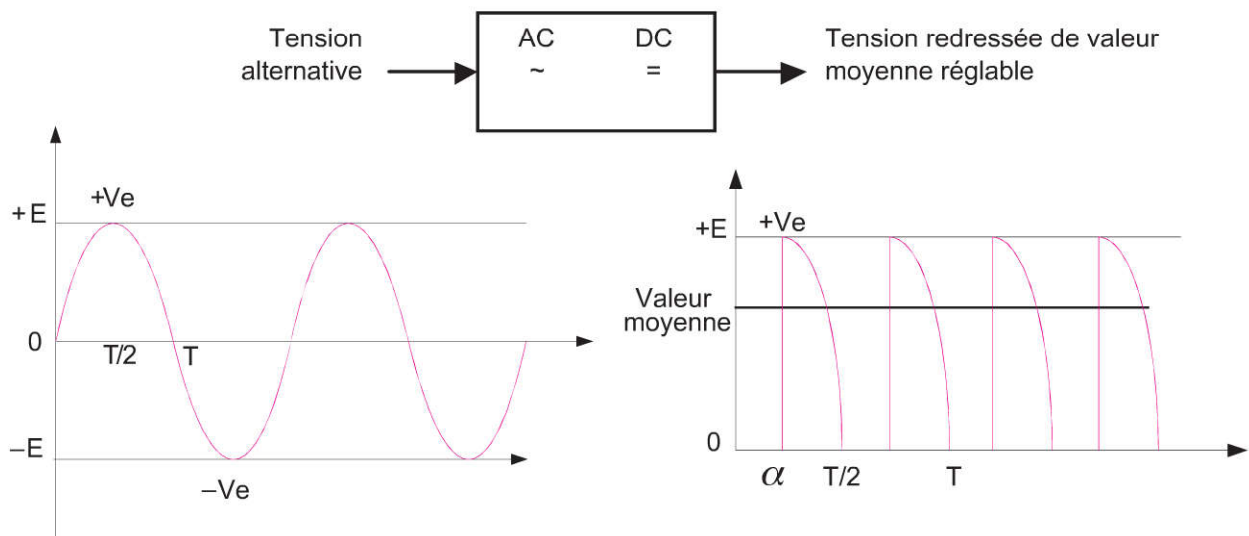


Remarque

Le transformateur est un convertisseur d'énergie électrique AC / AC isolé.

b) Convertisseur AC-DC

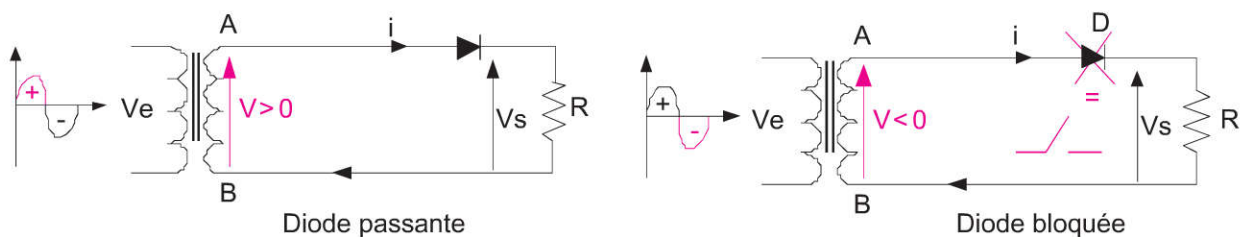
Un convertisseur AC-DC est un redresseur alternatif-continu qui existe sous différents types : redresseurs à diode, redresseurs à thyristors et redresseurs mixtes. Ci-dessous, on montre le procédé du convertisseur AC-AC :



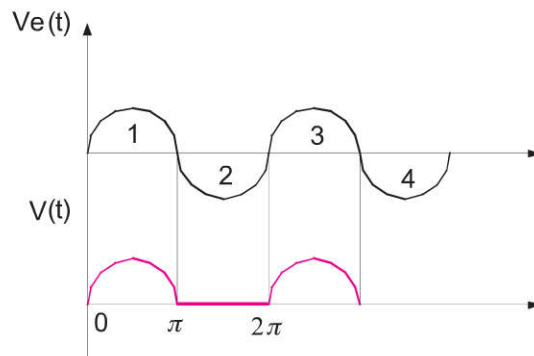
Dans un redresseur alternatif-continu monophasé, on peut étudier deux types de redressements qui sont de plus en plus utilisés :

- redressement mono alternance ;
- redressement double alternance.

Le **convertisseur AC/DC simple redressement** est construit autour d'une seule diode (fiche 31).



Dans les figures ci-dessus, quand la tension de $V_A > V_B$, la diode laisse passer le courant dans la charge R pendant une alternance ; par contre, quand la tension de $V_A < V_B$, la diode bloque le passage du courant. C'est seulement dans une alternance positive que la diode permet au courant de traverser.



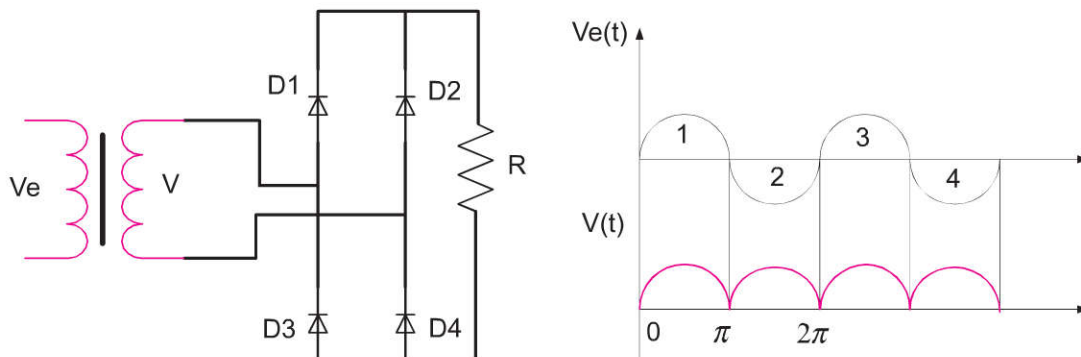
En négligeant la résistance de transformateur et la résistance de diode, le courant maximum sera

$$I_{\max} = \frac{V_{\max}}{R} \text{ et le courant moyen est : } I_{\text{moy}} = \frac{I_{\max}}{\pi}.$$

- La valeur efficace d'un courant sinusoïdal redressé à simple alternance est $I_{\text{eff}} = \frac{I_{\max}}{2}$.
- Les valeurs maximale et moyenne s'écrivent $V_{\max} = \sqrt{2} \cdot V_{\text{eff}}$ et $V_{\text{moy}} = \frac{V_m}{\pi} = \frac{R \cdot I_{\max}}{\pi}$.
- La valeur efficace de la tension de sortie est $V_{\text{eff}} = \frac{V_{\max}}{2}$.
- Le facteur de forme dans un redresseur à simple alternance sera $F = \frac{V_{\text{eff}}}{V_{\text{moy}}}$.
- Le rendement maximal de la conversion du courant alternatif en courant continu est alors :

$$\eta = \frac{P_{\text{continu}}}{P_{\text{alternatif}}} \text{ avec } P_{\text{continue}} = \frac{V_{\max}^2}{\pi^2 \cdot R^2} \text{ et } P_{\text{alternatif}} = \frac{V_{\max}^2}{4 \cdot R^2} \Rightarrow \eta = \frac{4}{\pi^2}$$

Le **redressement à double alternance** met en œuvre un pont de **quatre diodes**.



$$\text{La tension moyenne de sortie est } V_{\text{moy}} = \frac{2 \cdot \sqrt{2} \cdot V}{\pi}.$$

Si on installe les thyristors à la place des diodes dans la figure ci-dessus, la tension moyenne de

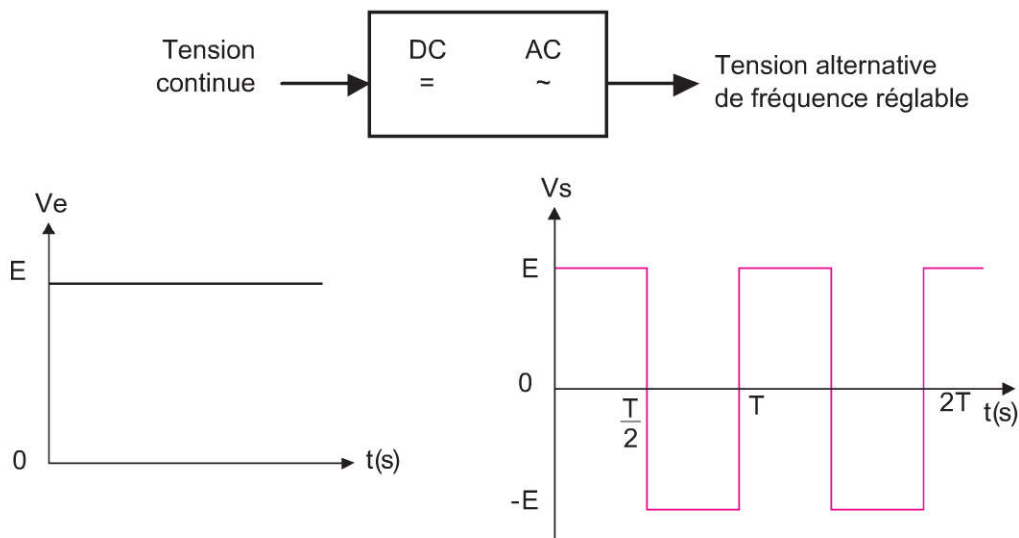
sortie sera $V_{moy} = \frac{2 \cdot \sqrt{2} V}{\pi} \cdot \cos(\theta)$.

L'angle θ peut varier entre 0 et π . Entre 0 et $\pi/2$, la tension moyenne de la charge est positive et entre $\pi/2$ et π la tension de la charge est négative.

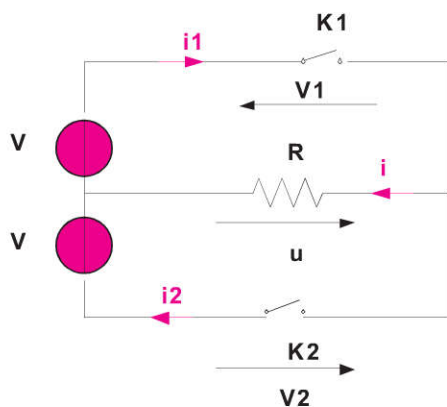
c) Convertisseur DC-AC

Un convertisseur DC-AC (continu-alternatif) s'appelle onduleur. Il peut assurer un transfert d'énergie entre une source de tension continue en entrée et une source de courant alternatif en sortie. Il existe deux types d'onduleurs : les onduleurs de tension et les onduleurs de courant.

- Les onduleurs de tension peuvent être alimentés par batterie, pile à combustible ou encore cellules photovoltaïques.
- Par contre, la source d'un onduleur de courant peut être un hacheur ou un redresseur en série avec une inductance. On utilise l'onduleur de courant seulement en haute puissance à l'aide de thyristors. Les nouveaux onduleurs actuels sont les onduleurs de tension.

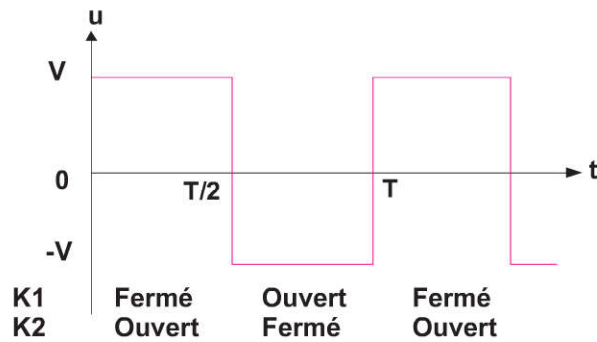


On explique ci-dessous le principe d'un onduleur de tension à deux interrupteurs avec une charge purement résistive. Il s'agit d'actionner alternativement les interrupteurs K1 et K2 durant des intervalles de temps réguliers.

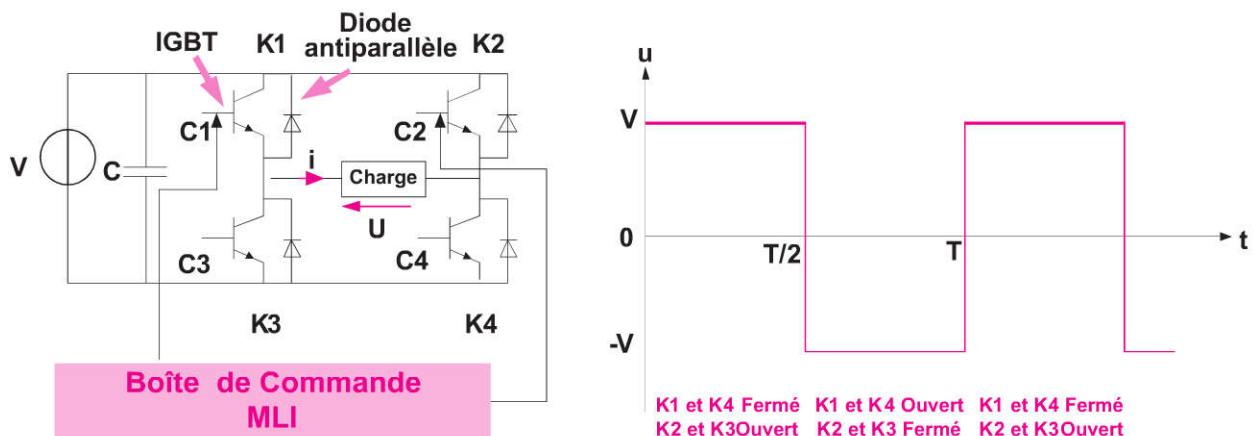


$$\begin{aligned} i &= i_1 - i_2 \\ V - v_1 - u &= 0 \\ V + u - v_2 &= 0 \\ i &= u / R \end{aligned}$$

La forme de tension en sortie est montrée ci-dessous :



Ci-dessous on trouve un onduleur de tension à quatre interrupteurs.



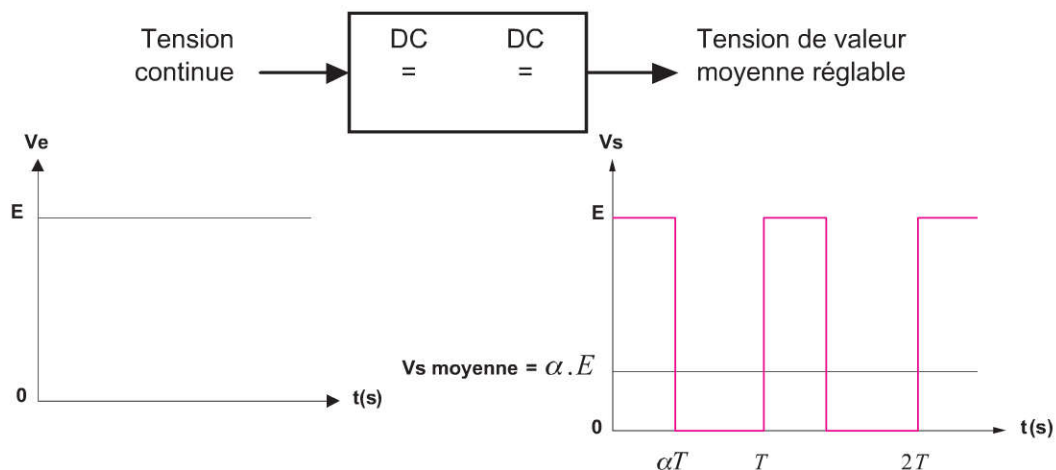
Dans cet onduleur, on utilise des IGBTs avec des diodes antiparallèles intégrées. Les commandes pour l'ouverture et la fermeture sont $C_3 = \overline{C_1}$ et $C_4 = \overline{C_2}$.

La tension instantanée aux bornes de la charge est $U = (C_m - C_n) \cdot V$ avec $1 \leq m, n \leq 4$ et $\overline{U} = G \cdot V_c$ est sa tension moyenne.

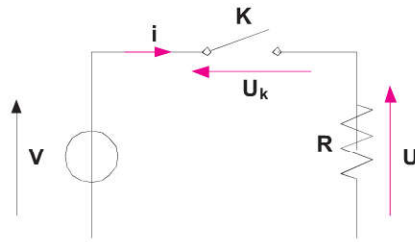
G est le gain et V_c est la tension de notre commande qui est utilisée pour créer quatre signaux pour les interrupteurs.

d) Convertisseur DC-DC

Un convertisseur DC-DC s'appelle un hacheur. Ce sont des convertisseurs d'énergie qui font transiter l'énergie électrique d'une source continue vers une autre source continue.



Le schéma montre le principe de fonctionnement d'un hacheur série (l'interrupteur K est monté en série entre la source et la charge) sur une charge résistive. L'interrupteur modélise un transistor.

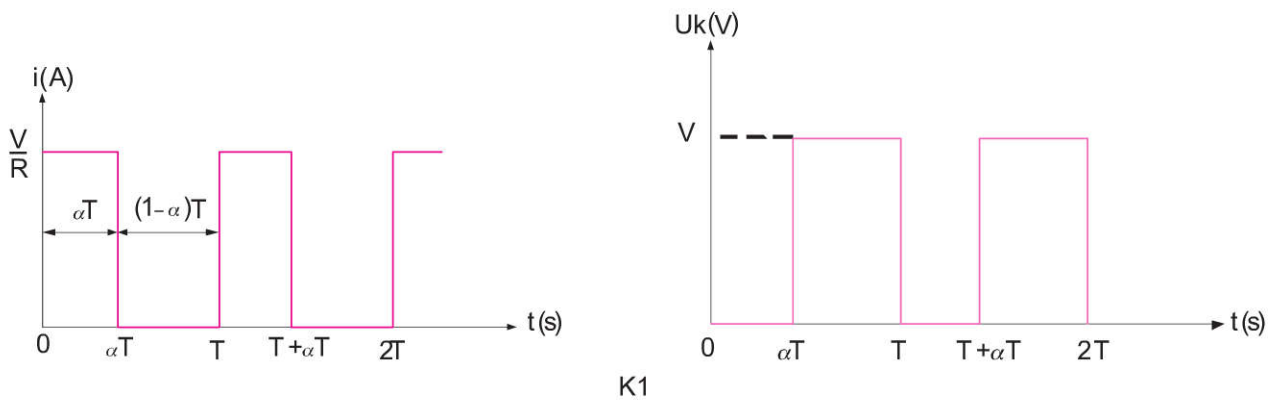


On choisit une période T et une fraction α de cette période. Le rapport cyclique α ($0 < \alpha < 1$, sans unité) gère la valeur moyenne de sortie U :

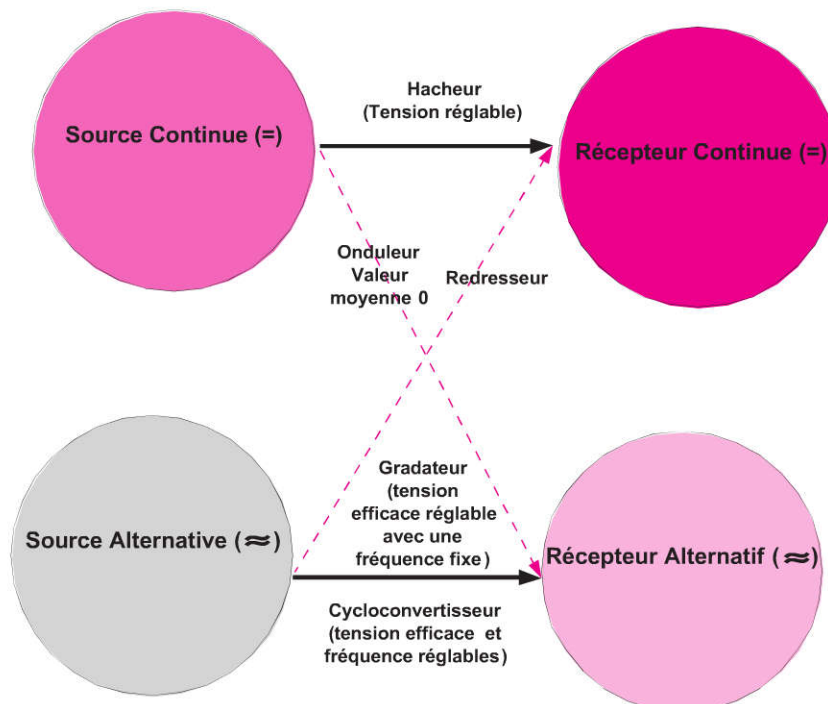
- de 0 à $\alpha \cdot T$: K est fermé et $U_K = 0$, $U = V$ et $i = U/R = V/R$;
- de $\alpha \cdot T$ à T : K est ouvert et $i = 0$, $v = Ri = 0$ et $u_K = V$.

Nous pouvons calculer la valeur moyenne de la tension en sortie du hacheur :

$$\overline{U} = \frac{\alpha \cdot T \cdot V + (1 - \alpha) \cdot T \cdot 0}{T} = \alpha \cdot V$$



3. EN PRATIQUE



73 Dimensionnement d'un dissipateur

Mots clés

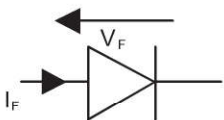
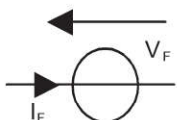
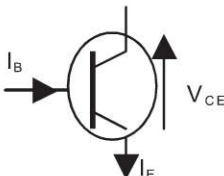
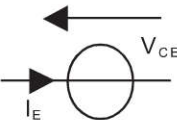
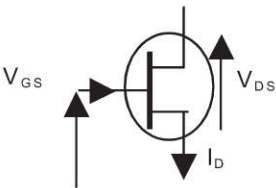
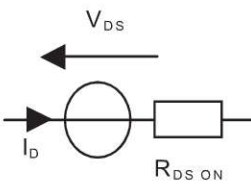
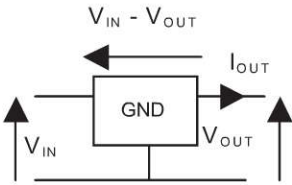
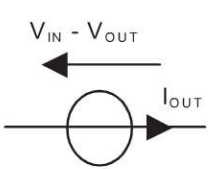
Puissance, commutation, conduction, résistance thermique, température de jonction.

1. EN QUELQUES MOTS

2. CE QU'IL FAUT RETENIR

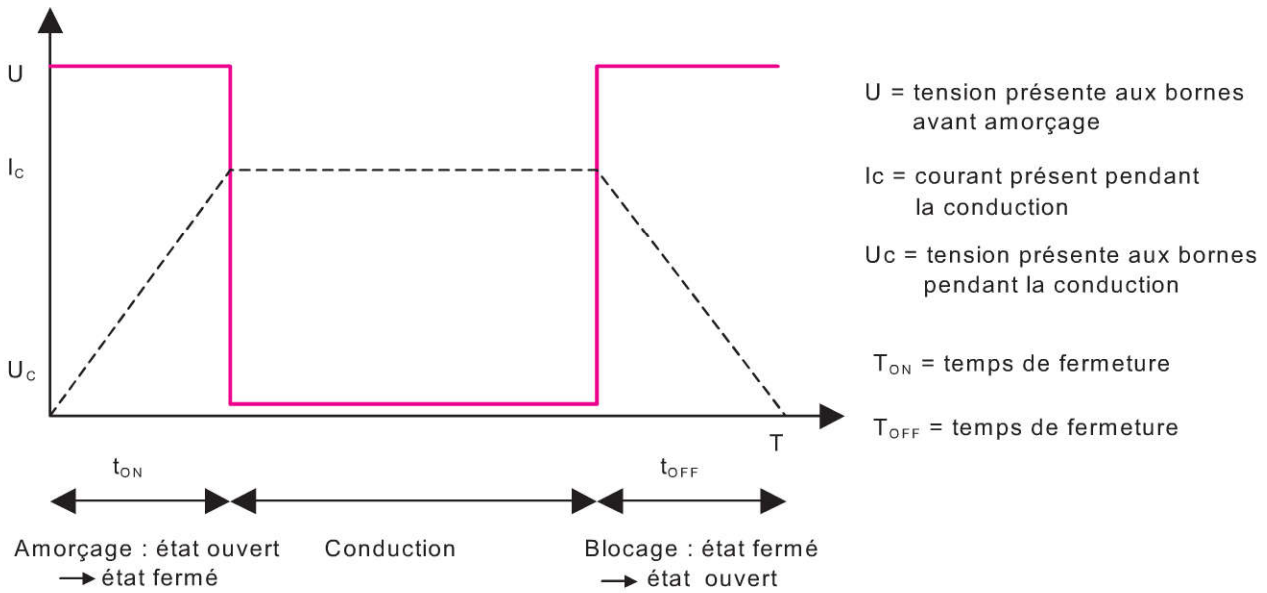
a) Puissance à dissiper

► Puissance dissipée lors de la conduction (P_C)

Composant	Symbole	Modèle en conduction	Puissance dissipée (valeurs moyennes de V et I sauf I_{Deff})
Diode			$P_C = V_F \times I_F$
Transistor bipolaire (NPN ou PNP)			$P_C = V_{CE} \times I_E$
Transistor MOS (canal P on N)			$P_C = V_{DS} \times I_D + R_{DS\ ON} \times I_{Deff}^2$
Régulateur de tension (78XX-79XX-LMXX...)			$P_C = (V_{IN} - V_{OUT}) \times I_{OUT}$

► Puissance dissipée lors de la commutation pour les semi-conducteurs (P_{COM})

Cette caractéristique ne prend pas en compte les particularités de chaque grande famille de semi-conducteurs mais donne une allure approchée qui est un des cas le plus défavorables (puissance à dissipée maximale).



Les pertes par commutation correspondent au produit $u(t) \propto i(t)$ durant t_{ON} et t_{OFF} en valeur absolue, ceci à chaque période de découpage T .

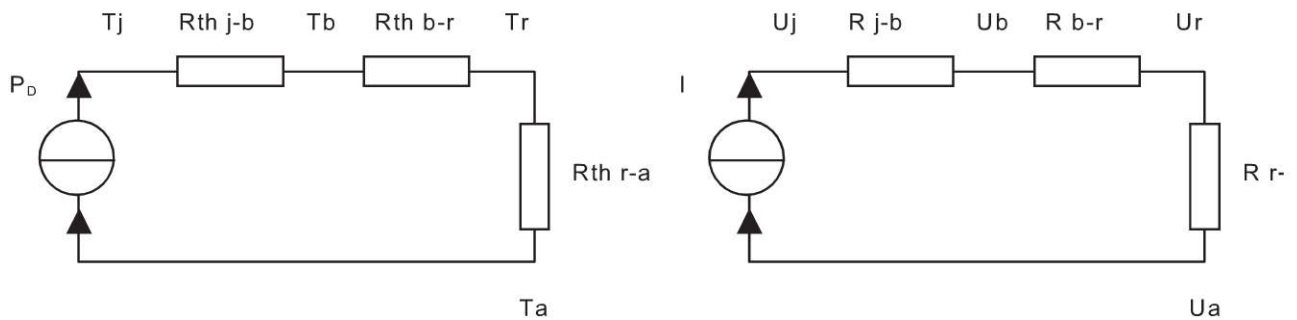
$$P_{COM} = \frac{w}{T} = w \times F = F \times \left[\int_0^{ton} u(t) \times i(t) dt + \int_0^{toff} u(t) \times i(t) dt \right]$$

Dans notre cas, on a $P_{COM} = \frac{1}{2} U_C \times I_C \times [t_{ON} + t_{OFF}] \times F$

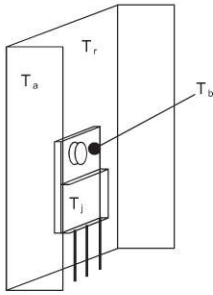
a puissance totale à dissiper : $P_D = P_C + P_{COM}$

b) Les différentes résistances thermiques

La puissance P_D est à dissiper depuis la jonction du semi-conducteur. En régime établi, le schéma thermique équivalent est à rapprocher de la loi d'ohm dans le domaine électrique :



$$\left\{ \begin{array}{l} P_D = \frac{T_j - T_b}{R_{th_{j-b}}} = \frac{T_b - T_r}{R_{th_{b-r}}} = \frac{T_r - T_a}{R_{th_{r-a}}} \\ R_{th_{j-b}} + R_{th_{b-r}} + R_{th_{r-a}} = (T_j - T_a) / P_D \end{array} \right\} \quad \left\{ \begin{array}{l} I = \frac{U_j - U_b}{R_{j-b}} = \frac{U_b - U_r}{R_{b-r}} = \frac{U_r - U_a}{R_{r-a}} \\ R_{j-b} + R_{b-r} + R_{r-a} = (U_j - U_a) / I \end{array} \right\}$$



La puissance devant être dissipée par la jonction doit traverser successivement le boîtier → la graisse de silicone ou le mica (afin de faciliter la transmission de la puissance et l'isolement électrique boîtier sur le radiateur) → le radiateur → extérieur.

On associe à chaque transfert une résistance thermique :

- Jonction → boîtier = $R_{th\ j-b}$ ou noté $j-c$ (jonction-case) ;
- Boîtier → radiateur = dépend du montage :
 - sans graisse silicone $R_{th\ b-r} \cong 0,3\ ^\circ C/W$;
 - avec graisse silicone $R_{th\ b-r} \cong 0,2\ ^\circ C/W$;
 - sans mica 100 μm sans graisse $R_{th\ b-r} \cong 1,5\ ^\circ C/W$

avec $R_{th\ b-r} \cong 0,6\ ^\circ C/W$;

- sans mica 50 μm sans graisse $R_{th\ b-r} \cong 1,25\ ^\circ C/W$ avec $R_{th\ b-r} \cong 0,4\ ^\circ C/W$.

Radiateur Ambiant = dépend du radiateur choisi. Ils sont très souvent teints en noir de façon à mieux rayonner la puissance à dissiper.

Dans le cas d'un montage horizontal, il est nécessaire de majorer la valeur de $R_{th\ r-a}$ de 20 %.

c) Choix d'un radiateur

Le critère de choix d'un radiateur ne tient qu'à la relation :

$$R_{th\ r-a} \text{ calculé} > R_{th\ r-a} \text{ du radiateur}$$

En partant de la relation du paragraphe précédent, on obtient :

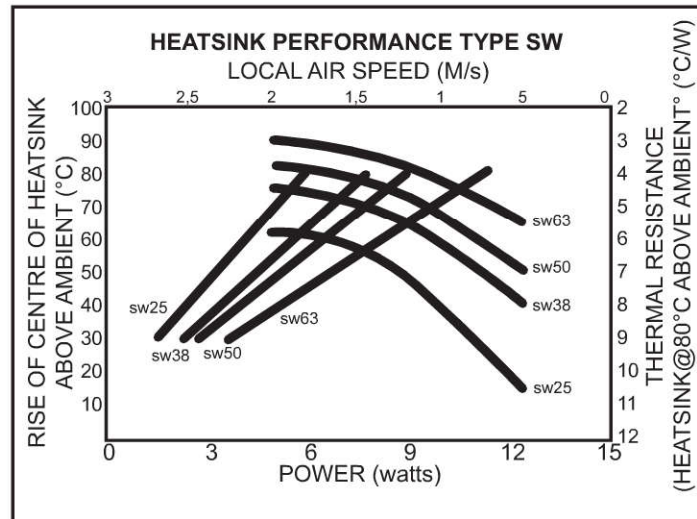
$$R_{th\ r-a} = \frac{T_{j\ MAX} - T_{a\ max}}{P_{D\ MAX}} - R_{th\ j-b} - R_{th\ b-r}$$

Un facteur de 1,5 peut être affecté à la valeur de $R_{th\ r-a}$ de façon à être certain de la bonne dissipation de la puissance. Il existe plusieurs façons d'exprimer les caractéristiques d'un radiateur en fonction des constructeurs.

► 1^{re} méthode

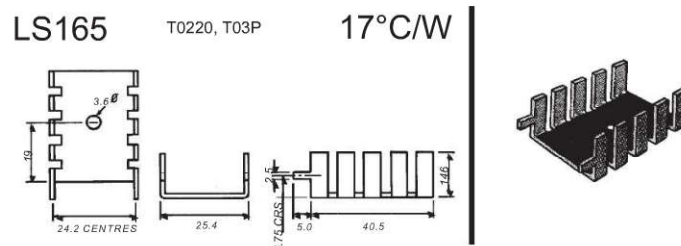
On choisit le radiateur en fonction de la puissance à dissiper et de l'écart de température entre le boîtier et l'ambiant. Dans ce cas, pas besoin de calculer $R_{th\ r-a}$. On obtient directement la référence du radiateur.

Le refroidissement est de type forcé (ventilateur), en connaissant la vitesse de l'air brassé et la valeur $R_{th\ r-a}$, on obtient la référence.



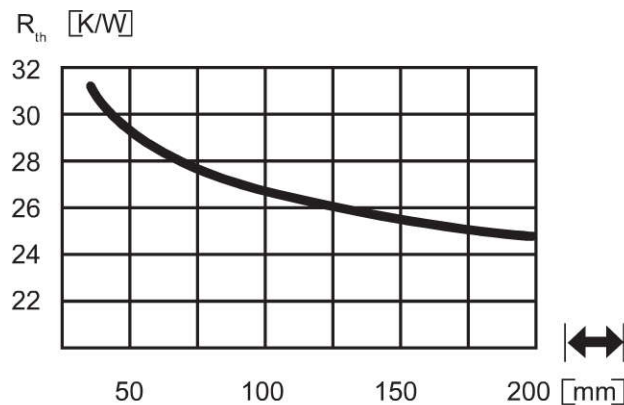
► 2^e méthode

Le constructeur donne juste la valeur de $R_{th\ r-a}$, les dimensions et le type de boîtier qui peut être monté.



► 3^e méthode

La valeur de $R_{th\ r-a}$ est donnée en °C/W ou en K/W (ce qui ne change rien car $1\ K = 1\ ^\circ C + 273,15$: il s'agit juste d'un décalage, une résistance thermique $1\ ^\circ C/W = 1\ K/W$) en fonction de la longueur de ce dernier s'il est peut-être découpé.



Remarque

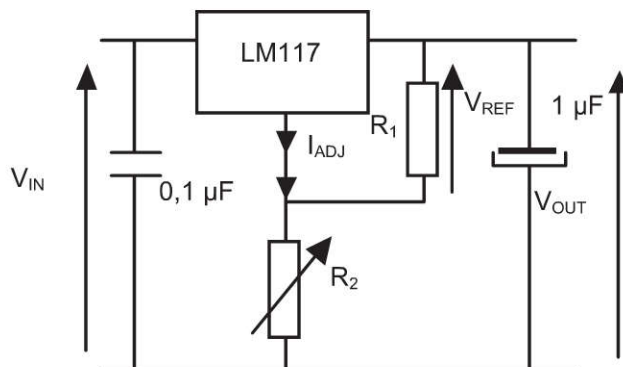
Dans le cas d'un montage avec plusieurs composants sur le même radiateur, il faut s'assurer que les boîtiers soient isolés les uns des autres (mica) afin de ne pas créer de court-circuit. En effet, les potentiels des boîtiers peuvent être différents.

Le calcul de $R_{th\ r-a}$ doit être redéfini par le schéma équivalent du départ. T_j sera imposée par la valeur la plus basse des composants montés ensemble.

3. EN PRATIQUE

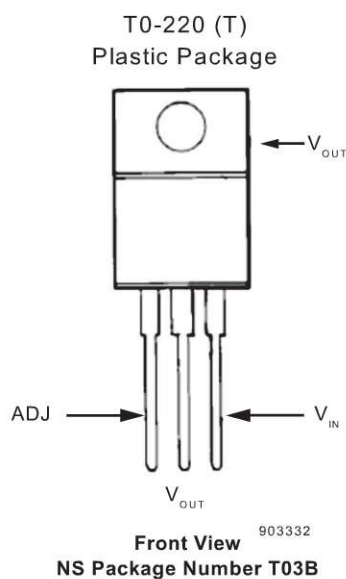
Un régulateur de type LM117 (package TO220 noté T03B dans la documentation constructeur) est utilisé pour obtenir une tension de sortie ajustable de 20 à 1,25 V par le réglage d'un potentiomètre. Le courant max de sortie sera fixé à 1 A sous une température ambiante max de 40 °C.

Schéma :



La fiche constructeur définit les valeurs suivantes :

- $V_{IN} < 28V$, $V_{REF} = 1,25 V$,
- $V_{OUT} = V_{REF} \times (1 + R_2/R_1) + I_{ADJ} \times R_2$, $I_{OUT} = 1,5 A$ max
- $R_1 = 1\ 200$ et $R_2 = 20\ k$
- $T_j = 150\ ^\circ C$
- j-c (notation constructeur) = $2\ ^\circ C/W$



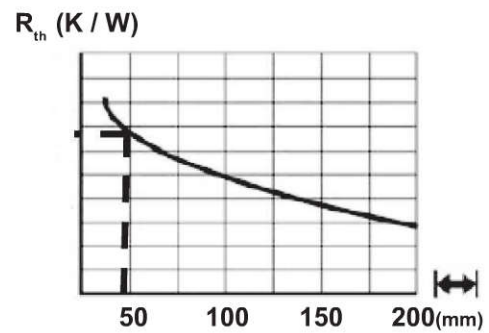
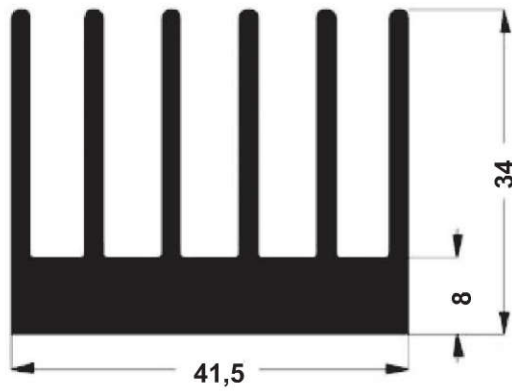
Pour déterminer la valeur $R_{th\ j-a}$, vous avez la possibilité de partir de la jonction elle-même ou déjà faire abstraction du montage du radiateur (utilisation de mica et de graisses).

La puissance dissipée dans le pire des cas correspond à une tension de 20 V en sortie et un courant de 1,2 A soit $P_D = (20 - 1,2) \times 1 = 17,8\ W$

En partant de la jonction :

$$R_{th\ r-a} = \frac{150 - 40}{17.8} - 2 - 0.6 = 3.6\ ^\circ C / W$$

Il faut que la résistance thermique du radiateur soit inférieure à $3,6\ ^\circ C/W$, ce qui est possible avec le radiateur ci-dessous pour une hauteur de 50 mm.



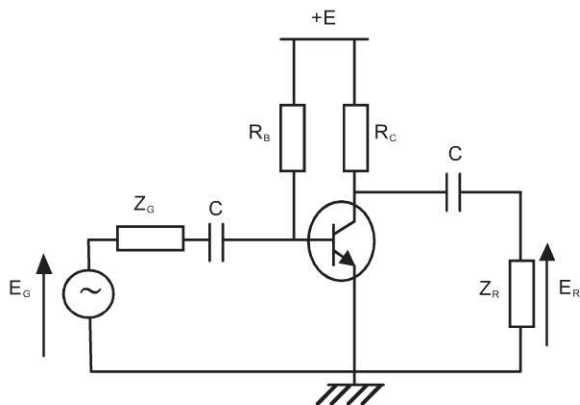
a) Pour conclure

La température de jonction impose la différentielle de potentiel thermique maximale entre elle-même et l'ambient.

Seule la valeur de la résistance thermique $R_{th_{r-a}}$ permet de choisir le dissipateur. Cette résistance thermique dépend de la longueur du dissipateur et de son type de convection (forcée ou naturelle).

74 Synthèse des trois montages à base de transistor bipolaire

Montages : émetteur, collecteur et base commune



G

$$\beta \times \frac{R_B}{R_S - r} \times \frac{R_{pRC}}{R_{pRC} + R_{ER}}$$

G_V

$$-\frac{Z_{EA}}{R_T} \times G_I \quad R_T = \rho // R_C // R_{ER}$$

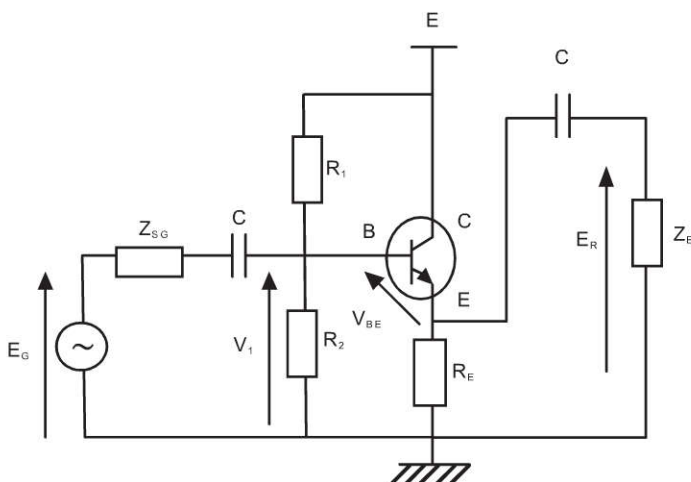
Z_{EA}

$$\frac{R_B \times r}{R_B + r}$$

Z_{SA}

$$R_{pRC} = \frac{1}{\frac{1}{\rho} + \frac{1}{R_C}}$$

Montages : émetteur, collecteur et base commune



G

$$R_B = \frac{1}{\frac{1}{R_1} + \frac{1}{R_2}}$$

$$-(\beta + 1) \times \frac{R_B}{R_S - Z'_E} \times \frac{R_E}{R_E + R_{ER}}$$

G_V

$$R_T = \frac{1}{\frac{1}{R_C} + \frac{1}{R_{ER}}}$$

$$\frac{(\beta + 1) \times R_T}{r + (\beta + 1) \times R_T} \cong 1$$

Z_{EA}

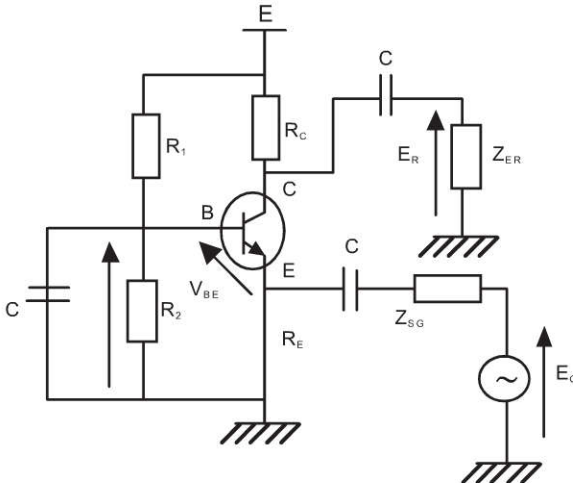
$$\frac{R_B \times Z'_E}{R_B + Z'_E}$$

$$Z'_e = r + (\beta + 1) \times R_T$$

Z_{SA}

$$\frac{R_{EQU} + R_E (\beta + 1)}{R_Z \times R_{EQU}}$$

$$R_{EQU} = \frac{1}{\frac{1}{R_G} + \frac{1}{R_B}} + r$$

Montages : émetteur, collecteur et base commune		
	G	$\frac{\beta \times R_C}{\left(\frac{r}{R_E} + \beta + 1\right) \times (R_C + R_{ER})}$
	G _V	$\beta \times \frac{R_T}{r}$ $R_T = \frac{1}{\frac{1}{R_C} + \frac{1}{R_{ER}}}$
	Z _{EA}	$\frac{r \times R_E}{r + (\beta + 1) \times R_E}$
	Z _{SA}	R_C

75 Rappels mathématiques

Mots clés

Intégration, fonctions trigonométriques, fonctions de Bessel, fonctions logarithmiques, fonctions exponentielles.

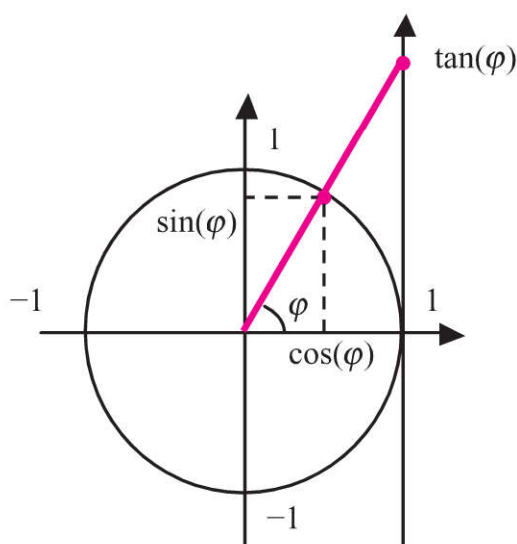
1. EN QUELQUES MOTS

Certains rappels mathématiques utiles pour les calculs présentés dans cet ouvrage sont donnés dans cette fiche.

2. CE QU'IL FAUT RETENIR

Propriétés des fonctions trigonométriques

Les fonctions sinus et cosinus sont définies par rapport au cercle trigonométrique qui est de rayon unité.



• Propriétés fondamentales

Théorème de Pythagore :

$$\sin^2(t) + \cos^2(t) = 1$$

Fonction tangente : $\tan(t) = \sin(t)/\cos(t)$

La fonction cosinus est paire :

$$\cos(t) = \cos(-t)$$

La fonction sinus est impaire :

$$\sin(t) = -\sin(-t)$$

• Valeurs remarquables

Angle	0°	30°	45°	60°	90°
	0 rad	$\pi/6$ rad	$\pi/4$ rad	$\pi/3$ rad	$\pi/2$ rad
sin	$\sqrt{0}/2 = 0$	$\sqrt{1}/2$	$\sqrt{2}/2$	$\sqrt{3}/2$	$\sqrt{4}/2 = 1$
cos	$\sqrt{4}/2 = 1$	$\sqrt{3}/2$	$\sqrt{2}/2$	$\sqrt{1}/2$	$\sqrt{0}/2 = 0$
tan	0	$\sqrt{3}/3$	1	$\sqrt{3}$	∞

• Formules d'addition

$$\begin{array}{l|l} \cos(a+b) = \cos(a) \cdot \cos(b) - \sin(a) \cdot \sin(b) & \sin(a+b) = \sin(a) \cdot \cos(b) + \sin(b) \cdot \cos(a) \\ \cos(a-b) = \cos(a) \cdot \cos(b) + \sin(a) \cdot \sin(b) & \sin(a-b) = \sin(a) \cdot \cos(b) - \sin(b) \cdot \cos(a) \end{array}$$

À partir des relations ci-dessus, nous pouvons écrire :

$$\begin{aligned}\cos(2a) &= \cos(2 \cdot a) - \sin(2 \cdot a) = 1 - 2 \cdot \sin(2 \cdot a) = 2 \cdot \cos(2 \cdot a) - 1 \\ \sin(2 \cdot a) &= 2 \cdot \sin(a) \cdot \cos(a)\end{aligned}$$

On en déduit : $\cos(a) = 1 - 2 \cdot \sin(a/2)$

$$\sin(a) = 2 \cdot \sin(a/2) \cdot \cos(a/2)$$

- Remplacement par rapport à $\pi/2$ $\cos(\pi/2 - t) = \sin(t)$ $\sin(\pi/2 - t) = \cos(t)$
 $\cos(\pi/2 + t) = -\sin(t)$ $\sin(\pi/2 + t) = \cos(t)$

On dit que la fonction cosinus est en avance de phase de $\pi/2$ par rapport à la fonction sinus.

- Remplacement par rapport à π $\cos(\pi - t) = -\cos(t)$ $\sin(\pi - t) = \sin(t)$
 $\cos(\pi + t) = -\cos(t)$ $\sin(\pi + t) = -\sin(t)$

Propriétés des fonctions logarithmiques

- Associativité $\log(a \cdot b) = \log(a) + \log(b)$
 $\log(a/b) = \log(a) - \log(b)$
 $\log(a^n) = n \cdot \log(a)$

- Réciprocité $e^{\ln(x)} = x$ $\exp(\log(x)) = x$
 $\ln(e^x) = x$ $\log(\exp(x)) = x$

Fonctions de Bessel de première espèce

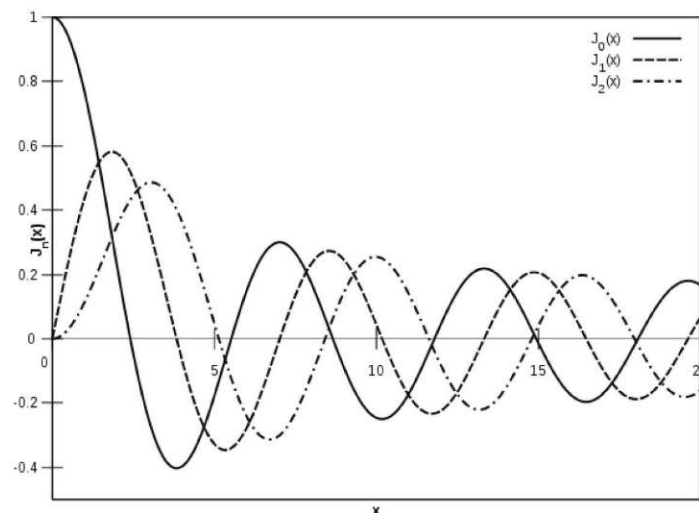
Les fonctions de Bessel sont des solutions y de l'équation différentielle de Bessel :

$$x^2 \cdot \frac{d^2 y}{dx^2} + x \cdot \frac{dy}{dx} + (x^2 - n^2) \cdot y = 0$$

Les fonctions de Bessel de première espèce J_n sont les solutions de l'équation différentielle ci-dessus qui sont définies en 0.

Elles sont définies par : $J_n(x) = \left(\frac{x}{2}\right)^n \times \sum_{p=0}^{\infty} \left(\frac{(-1)^p}{2^{2p} \cdot p! \cdot (n+p)!} \cdot x^{2p} \right)$

Voici le tracé des trois premières fonctions de Bessel de première espèce :



Index alphabétique

A

actionneur 194
adaptation 54
 d'impédance 144
 en courant 54
 en puissance 54
 en tension 54
additivité des gains 50
admittance 18
alternatif 112
alternée 128
ampère 4, 36
amplificateur différentiel 152
amplification 112
amplitude 24, 124
analyse 38, 42
analyse des circuits 46
AOP 54
 idéal 150
 réel 150
apport cyclique 226
architecture 182
ASIC 174
asservissement 204
atome 2
atténuation 70, 198

B

balayage
 horizontal 24
 vertical 24
banc de test 186
bande
 d'énergie 2
 de base 208
 de conduction 2
 de valence 2
bascule
 D 172
 JK 172
 RS 172
base 92

BCD 166
binaire 162
 décalé 166
bloc logiques configurables 178
branche 28
bruit
 blanc 216
 de grenaille 216
 rose 216
 thermique 216
bus 191

C

calibre 24
CAN 194
 à approximations successives 200
canal 132, 206
capacité 16, 138
capteur 194
caractéristique
 d'entrée 96, 136
 de sortie 96, 136
 de transfert 96
causalité 48
cellule
 de Rauch 156
 de Sallen-Key 156
champ électrique 84
charge électrique 2, 8
circuit
 intégré 170
 logiciels et matériels 174
 RLC série 74
CLB 178
CNA 194
collecteur 92
 commun 54
commutation 92, 232
comparateur de phase 204
complément à 2 166
composants passifs 62
compteur 172

condensateur 12
 conductance 18, 107
 conduction 78, 128, 232
 contrôle 191
 convolution 60
 coulomb 4, 36
 courant
 alternatif 20, 30
 de gâchette nul 136
 de polarisation 150
 CPLD 174
 cycle de développement 174
 cycloconvertisseurs 226

D

dB 50
 dBm 50
 dBW 50
 DCB 166
 de Sauty 68
 découpler 112
 dénormalisation 158
 dépendant b 102
 déphasage 24
 description matérielle 182
 diagramme
 de Bode 50, 70, 198
 de Fresnel 34
 de Nyquist 60
 dirac 42
 dispersion 136
 dissipée 124
 diviseur
 de courant 62
 de fréquence 204
 de tension 62
 domaine fréquentiel 38
 donnée 191
 dopage 80
 droite de charge et d'attaque 96
 dynamique 124

E

échantillonnage 194
 effet transistor 92
 électron 78
 émetteur 92, 206
 end 182

énergie 10, 218, 226
 ENOB 200
 entity 182
 équation
 différentielles 46
 logique 162

F

facteur
 d'amplification 107, 138
 de puissance 32
 de qualité 32, 74
 faible gain inversé 140
 farad 12
 filtrage 50, 204
 filtre 74
 actif 152
 passe-bas 70, 198
 passe-haut 70, 198
 flash 200
 fonction
 de Bessel 240
 échelon 42
 exponentielle 34, 240
 logarithmique 240
 porte 42
 trigonométrique 240
 fondamentale 38
 force 8
 électromotrice 10
 fournie 124, 128
 FPGA 174
 fréquence 20, 24
 de Nyquist 194
 de résonance 74

G

gabarit 158
 gain 114, 144
 en boucle ouverte 150
 en courant unitaire 122
 en tension important 122
 en tension unitaire 118
 galvanomètre 68
 générateur
 équivalent 66
 idéal 10
 réel 10

germanium 78
 gradateurs 226
 GTO 222

H

hacheurs 226
 harmonique 38
 henry 12
 hertz 20
 hexadécimal 166
 horloge 172

I

IGBT 222
 IGCT 222
 impédance 107, 138
 d'entrée élevée 144
 d'entrée et de sortie 114
 d'entrée grande 118, 122, 140
 de sortie faible 118, 122, 144
 équivalente 66
 opérationnelle 46
 indépendant
 b 102
 indice de modulation 208
 inductance 12, 16
 influence, récepteur 114
 instanciation 186
 instruction 191
 concourante 186
 séquentielle 186
 intégrateur 152
 intégration 240
 interconnexion 178
 intermédiaire 140
 interrupteur 218, 220
 inverseur 152
 isolants 78

J

Jonction 92
 joule 4, 32, 36
 JTAG 178

L

latch 172
 LED 88
 linéaire 92

liste de sensibilité 186
 logique booléenne 170
 loi
 d'Ohm 10
 de Coulomb 8
 de Kirchhoff 66
 de nœuds 28
 des mailles 28

M

matrice
 admittance 56
 chaîne 56
 impédance 56
max slew rate 150
 métaux 78
 microcontrôleurs 191
 mise en cascade 56
 mode 146
 différentiel 146
 modèle 107, 138
 équivalent petits signaux 146
 modulation
 BPSK 212
 CPFSK 212
 d'amplitude 208
 de fréquence 208
 de phase 208
 démodulation 206
 FSK 212
 MSK 212
 OOK 212
 QAM-N 212
 QPSK 212
 modulé 208
 montages non inverseur 152
 MOSFET 222
 mouvement brownien 216

N

normalisation 158
 notation complexe 30
 numérique 174
 numérisation 194

O

octal 166
 ohm 12, 54

ondes incidentes 18
onduleurs 226
opérateur
 ET 162
 NAND 162
 NON 162
 NOR 162
 OU 162
 XOR 162
orbite 2

P

PAL 174
paramètres 107, 138
pas de quantification 200
peigne de Dirac 48
perdue 128
période 20
phase 20
photodiodes 88
PIN 88
pincement 132
pipeline 200
plan mémoire 191
PLL 204
point de fonctionnement 96, 102
polarisation 112
 directe et inverse 84
pont
 de diode 226
 de mesure 68
 de Wheatstone 68
process 182, 186
processeur 191
puissance 50, 124, 128, 218, 232
 active 32
 moyenne 32
 réactive 32
pulsation réduite 74

Q

quadripôle 50
quantification 194

R

réactance 18
récepteur 206
redressement simple alternance 88

redresseurs 226
réfléchies 18
régime sinusoïdal 30
réponse impulsionnelle 48, 60
réseau 28
résistance 12, 16, 18, 30
 interne 10
 thermique 232
résistivité 78
résolution 200
ressources dédiées 178

S

schéma équivalent 220
sélectivité 158
semi-conducteur 2, 220
 extrinsèque 84
 intrinsèque 78
 type N 80
 type P 80
seuil 132
SFDR 200
shottky 88
sigma-delta 200
signal 112
 modulant 208
 périodique 38
signaux sinusoïdaux 20
signe
 VGS 132
silicium 78
SNR 200, 216
sommateur 152
source 206, 218, 226
spectre 42, 158
std_logic 182
suiveur 152
susceptance 18
symétrie 158
 hermitienne 60
synthèse 38, 42
système
 international 4, 36
 linéaire 48

T

tableau de Karnaugh 162
température de jonction 232

tension différentielle 146

théorème

de Boucherot 32

de De Morgan 162, 170

de Millman 68

de Shannon 194

thyristor 222

tolérance 16

transformation 218, 226

transformée de Fourier 46

transistor bipolaire 146, 222

transmise 18, 124

transmittance 74, 198

TRIAC 222

TRMC 146

trou 78

U

unité 4, 36

utile 128

V

VA 32

valeur

efficace 20

moyenne 20

normalisée 16

variable logique 162

VCO 204

vecteurs 8

volt 4, 36

W

watt 4, 32, 36

Z

zener 88

zone

active 132

ohmique 132